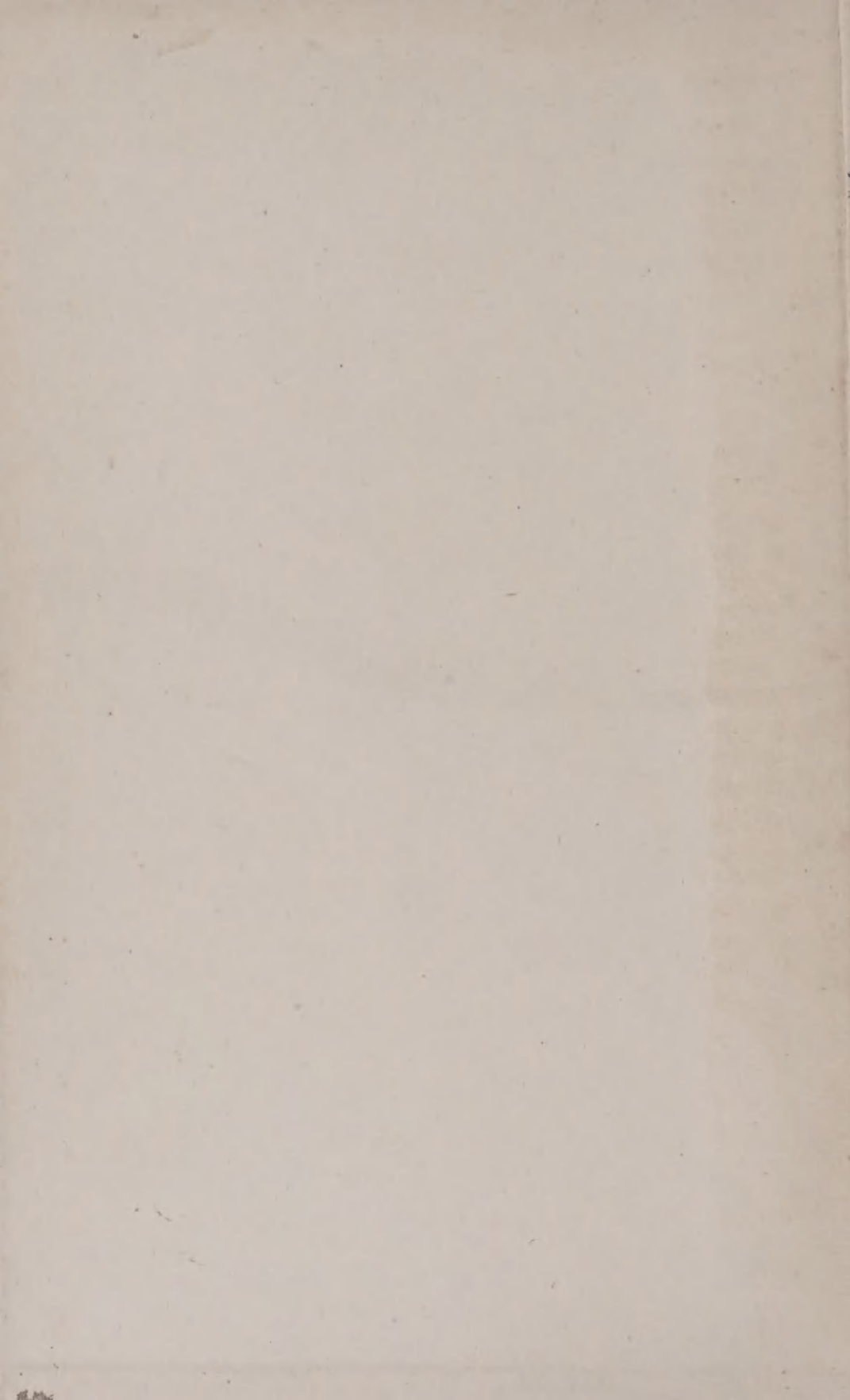


CFTRI-MYSORE



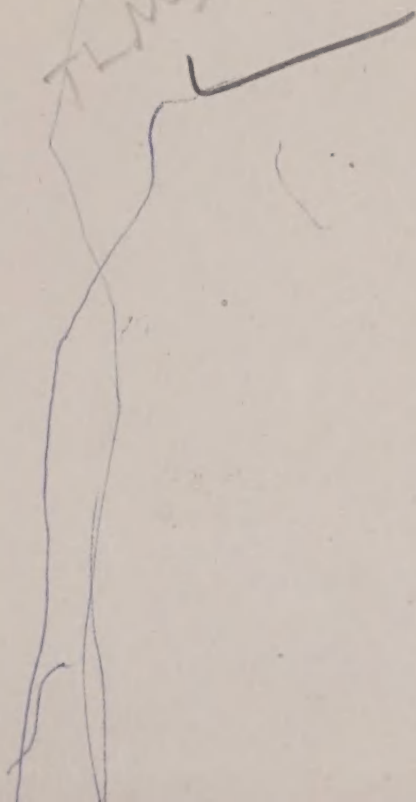
229

Chemical computa



1

TLMS





**CHEMICAL COMPUTATIONS
AND ERRORS**

BOOKS BY JOHN H. YOE

A LABORATORY MANUAL OF QUALITATIVE ANALYSIS
219 pages. 6 by 9. 7 figures. Cloth.

CHEMICAL PRINCIPLES

With particular application to qualitative analysis. 311
pages. 6 by 9. 29 figures, 2 colored plates. Cloth.

PHOTOMETRIC CHEMICAL ANALYSIS

Volume I. Colorimetry. 771 pages. 6 by 9. 72 figures.
Cloth.

Volume II. Nephelometry. 337 pages. 6 by 9. 44
figures. Cloth.

BY THOMAS B. CRUMPLER AND JOHN H. YOE

CHEMICAL COMPUTATIONS AND ERRORS

247 pages. 6 by 9. 15 figures. Cloth.

BY JOHN H. YOE AND LANDON A. SARVER

ORGANIC ANALYTICAL REAGENTS

339 pages. 6 by 9. 2 figures. Cloth.

CHEMICAL COMPUTATIONS AND ERRORS

BY

THOMAS B. CRUMPLER, PH.D.

*Professor of Chemistry and Head of the Department of Chemistry
The Tulane University of Louisiana*

AND

JOHN H. YOE, PH.D.

*Professor of Chemistry
University of Virginia*



NEW YORK

JOHN WILEY & SONS, INC.

LONDON: CHAPMAN & HALL, LIMITED

229 ✓

COPYRIGHT, 1940,

BY

THOMAS B. CRUMPLER AND JOHN H. YOE

All Rights Reserved

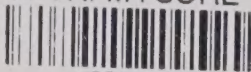
*This book or any part thereof must not
be reproduced in any form without
the written permission of the publisher.*

THIRD PRINTING, APRIL, 1947

E:b2

N47

CFTRI-MYSORE



229

Chemical computa.

PRINTED IN U. S. A.

Absolute certainty is a privilege of uneducated minds—and fanatics. It is for scientific folk an unattainable ideal. —C. J. KEYSER.

PREFACE

Since the advent of modern theoretical chemistry, based upon equilibrium and thermodynamics, it has become increasingly apparent that a chemist must have at least a working knowledge of mathematics. Not only must he be able to employ mathematical formulations and to conduct numerical operations, but also he must have the ability to interpret numerical results. The correlation of computation with interpretation appears to lie in a hazy domain between the functions of the teachers of mathematics and those of the teachers of the natural sciences. As a result, many students in elementary chemistry courses can solve with alacrity an abstract equation like

$$2x - 1 = 5$$

but when confronted with an equation such as

$$0.1 \left(\frac{\alpha^2}{1 - \alpha} \right) = 3.7 \times 10^{-5}$$

they often find both the solution and the interpretation difficult or impossible. How many figures should be retained in the answer? Of what importance is the source of the data in the problem; who determined them, and how? These and similar questions arise.

This book is intended for undergraduate students specializing in chemistry. It may be used either as a textbook for a course on chemical computations and errors or as a supplementary textbook in any one of several courses in a regular chemical curriculum, e.g., quantitative analysis or physical chemistry. In fact, use of the book might extend over a period of two or three years, the material in the early chapters being taken up in the student's second year and the later (and more advanced) chapters introduced in the third and fourth years. Furthermore, many graduate students in chemistry might profit by a study of much that is presented in this book. Indeed, it is quite possible that its contents (or the equivalent) might be considered a desirable minimum requirement in mathematics for graduate degrees in chemistry.

It is hoped that the book will also prove useful for home study to many chemistry students who do not have an opportunity to take formal courses on the topics presented, or who will use it as a reference book to review and to correlate material covered in their courses in mathematics, chemistry, and physics. With this in view, many illustrative examples are worked out in detail and the descriptive sections have been made as full as possible in a book of this scope. For instance, Chapter VII, "Classification of Errors," is an attempt to bring together and to discuss representative types of errors, their sources, magnitudes, and corrections that are encountered in the ordinary methods of making quantitative chemical measurements. Even such elementary topics as exponential numbers, logs, the slide rule, significant figures, and rules for computation are treated in the introductory chapters for the benefit of those students who need further drill on these subjects.

The book is not intended to replace the advanced textbooks and treatises on the various subjects presented. For those who want to extend their studies on certain phases of the work, a select and classified bibliography will be found in the Appendix. We have striven to produce a discussion of measurement and errors presented from the viewpoint of the chemist rather than that of a statistical mathematician. Such a presentation has necessitated a sacrifice of some mathematical rigor. Believing that many students of chemistry are graduated with an improper appreciation of the subject of measurement we have attempted an interpretation, not hesitating to point out some of the difficulties which beset quantitative measurements. We feel that these difficulties are often so little emphasized as to lead to false impressions and over-simplification in the minds of students.

Any claim to merit for this book will be based chiefly on the selection of material for the illustrative examples and exercises. Wherever possible, data of chemical measurements have been employed to illustrate both the mathematical methods and the applications of the principles of elementary theoretical chemistry. References to original sources are given in an attempt to help familiarize the student with the literature of chemistry.

The recommendations of the National Bureau of Standards¹ in regard to the spelling and abbreviation of units have been

¹ *Miscellaneous Publication M121*, 1936.

adopted, with one exception. We have not substituted cm^3 for *cc*, since we consider the latter to be simpler and so universally used as not to warrant a change.

Acknowledgment is made to the following companies for permission to use the quotations at the chapter headings indicated:

Chapter II, Comte (translated by Gillespie): *Philosophy of Mathematics*, p. 36, Harper and Brothers, New York, 1858.

Chapters IV and XI, Campbell: *Measurement and Calculation*, p. 1 and p. 186, Longmans, Green and Co., London, 1928.

Chapter V, Lewis and Randall: *Thermodynamics*, p. 8, McGraw-Hill Book Co., New York, 1923.

Chapter VI, Sullivan: *Gallio*, p. 57, E. P. Dutton and Co., New York, 1928.

Chapter VIII, Whitehead: *The Organization of Thought*, p. 15, J. B. Lippincott and Co., Philadelphia, 1917.

The quotation in the frontispiece is from Keyser: "Mathematical Philosophy," p. 120, E. P. Dutton and Co., New York, 1922. The electro for Figure 2 was furnished by Keuffel and Esser Co., and permission to use Figure 8 was granted by the Clark University Press. This cooperation is appreciated.

Our thanks are due Professor F. L. Brown, School of Physics, University of Virginia, and Dr. J. C. Elgin, professor of chemical engineering, Princeton University, who read the completed manuscript and made helpful suggestions. Professor W. S. Weedon, School of Philosophy, University of Virginia, read the manuscript for Chapters VI to XI and made valuable criticisms for which we are very grateful.

To Professor E. J. Oglesby, professor of engineering mathematics, University of Virginia, we owe a special acknowledgment of appreciation, not only for reading the entire manuscript and offering invaluable suggestions, but also for his inspiring lectures which served as a model for the presentation of much of the material in this book.

We are indebted to Mr. E. G. Ballard, Department of English, Virginia Military Institute, for his rendition of a quotation from *The Pearl*, and to Mr. K. L. McIntosh, Department of Chemistry, Tulane University, whose unselfish labor in connection with the preparation of the table of answers has contributed greatly towards insuring its accuracy.

Finally, we must thank Mrs. Margaret Lewis Crumpler and Mr. V. D. Napoleon for their assistance in reading the proofs and Mrs. Françoise Cheely Yoe for her painstaking care in typing the manuscript.

Although much effort has been expended to have the book free of mistakes, it is recognized that some may be found. We shall be obliged to anyone who may call attention to errors or to statements that should be revised. Criticisms and suggestions will be welcomed.

T. B. C.

J. H. Y.

NEW ORLEANS, LA.

CHARLOTTESVILLE, VA.

January, 1940

CONTENTS

CHAPTER	PAGE
I. COMPUTATION METHODS	1
Exponential Numbers	
Multiplication and Division—Powers and Roots—Addition and Subtraction.	
Logs (Logarithms)	
Historical—Multiplication by Logs—Division by Logs—Operations Involving both Multiplication and Division—Evolution—Involution—Negative Logs—Natural Logs.	
Slide Rule	
Historical—Multiplication—Division—Powers and Roots—Logs.	
II. SIGNIFICANT FIGURES	30
Definitions—Rules for Computation—Complex Cases.	
III. ALGEBRAIC SOLUTION OF NUMERICAL EQUATIONS	36
Linear Equations	
Stoichiometrical Computations—Simultaneous Linear Equations—Elimination.	
Determinants	
Solving Simple Simultaneous Equations—Determinants of Higher Order—Simultaneous Stoichiometrical Equations.	
Quadratic Equations	
Equations of Higher Degree	
IV. APPROXIMATE METHODS OF SOLVING EQUATIONS	53
Neglect of Added or Subtracted Quantity—Successive Approximations — Graphical Approximation — Newton's Method.	
V. INTERPOLATION AND EXTRAPOLATION	64
Linear Interpolation	
Rule of Proportional Parts—Inverse Interpolation.	
Non-Linear Interpolation	
Polynomials—Differences—Operators—Difference Formula of Newton—Variable Increments—Newton's Divided Difference Formula—Lagrange Formula.	
Inverse Non-Linear Interpolation	
Extrapolation	
Graphical Extrapolation—Extrapolation with Interpolation Formulas—Other Analytical Methods.	

CHAPTER	PAGE
VI. THEORY OF MEASUREMENTS	82
Standards	
Fundamental Standards	
Length—Mass—Time—Temperature—Historical—Area—	
Volume.	
Derived Standards	
Density—Specific Gravity—Energy—Pressure—Heat Quantity.	
Relative Standards	
Atomic Weights—Electrode Potentials—Activity.	
Measurement	
Observations	
Averages	
Errors	
Precision	
Accuracy	
VII. CLASSIFICATION OF ERRORS	96
Corrigible Errors	
Temperature	
Thermometric Calibration—Emergent Stem Correction—	
Glass Fatigue—Electrical Measurement.	
Pressure	
Temperature Correction for Barometer—Capillary Depres-	
sion—Correction for Gravity.	
Weighing	
Weight Calibration—Balance Arm Inequality—Air Buoy-	
ancy—Gravimetric Analysis—Incomplete Volatilization—	
Partial Decomposition—Crucible Loss—Adsorbed Mois-	
ture—Electrification—Solubility of Precipitate—Colloidal	
Dispersion—Contamination of Precipitate.	
Volume	
Calibration of Volumetric Ware—Standard Laboratory	
Temperature—Parallax—Liquid Films—Cleanliness—Out-	
flow Time—Viscosity of Liquid—Temperature Corrections.	
Congruity in Measurements	
Random Errors	
Residual Errors—Instrumental Errors—Personal Errors.	
VIII. STATISTICAL METHODS	125
Averages	
Mode—Median—Arithmetic Mean—Root Mean Square—	
Weighted Mean.	
Deviations	
Definition—Properties—Computation of Arithmetic Mean.	

Dispersion	
Range—Average Deviation—Standard Deviation—Computation of Standard Deviation.	
Distribution	
Normal Distribution—Class Interval.	

IX. THEORY OF ERRORS 141

Distribution of Errors	
Probability Theorems	
Historical—Mathematical Probability—Equally Likely Events—Complex Events—Permutations—Combinations—Limitations.	
Coin Tossing	
One Coin—Two Coins—Three Coins—Four Coins—Binomial Expansion—Frequency Distribution—Binomial Distribution.	
Curve of Error	
Derivation—Properties—Evaluation of Constant—Principle of Least Squares—Probability Integral—Probable Error—Verification of the “Law” of Errors.	

X. STATISTICAL INTERPRETATION OF MEASUREMENTS 174

Probable Error of Computed Result	
Sum or Difference—Product—Quotient—Powers or Roots—Logs—Exponential Functions—Complex Functions—Obtaining Congruity in Measurement.	
Probable Error of Arithmetic Mean	
Statistically Reliable Difference—Interpretation.	
Rejection of Observations	
Statistically Weighted Average	
Limited Number of Observations	
Precision Measures	
Absolute—Percentage—Parts per Thousand.	
Significant Figures	
Some Approximations	

XI. CURVE FITTING 202

Choice of Equation	
Finite Differences—Graphical Method—Smoothing of Data—Multiple Equations—Substituted Variables.	
Evaluation of Constants	
Interpolation Formulas—Graphical—Simultaneous Equations—Method of Averages—Method of Least Squares—“Goodness” of Fit.	

APPENDIX

	PAGE
BIBLIOGRAPHY	227
TABLES	230
Logs (Four-Place)—Error Integral—Atomic Weights—Answers to Problems	
AUTHOR INDEX	239
SUBJECT INDEX	243

CHEMICAL COMPUTATIONS AND ERRORS

CHAPTER I

COMPUTATION METHODS

When you can measure what you are speaking about and express it in numbers, you know something about it, and when you cannot measure it, when you cannot express it in numbers, your knowledge is of a meagre and unsatisfactory kind. It may be the beginning of knowledge, but you have scarcely in your thought advanced to the stage of a science.

—LORD KELVIN.

A physical quantity is measured in terms of an arbitrarily chosen standard unit. The measurement is expressed as the ratio of the magnitude of the quantity to the magnitude of the standard unit. It is interesting and often highly entertaining to choose at random several commonly used units of measurement and trace their formulation, as well as the etymology of their names. (See Chapter VI for some examples.) It is only in this way that one can grasp the rôle that fate or chance has played in fixing for us our scientific heritage that must serve as the unconsciously perceived background of all scientific thought and action.

The conception of number to express a ratio is very old compared with our civilization. The system of numerical symbols which we use was taken from the Arabs. The early Arabic mathematicians had formulated rules for combining these numerical symbols to obtain quickly and easily other ratios that could have been apprehended by the senses through the tedious processes of counting with the fingers and toes. Meanwhile, our less sophisticated ancestors were fumbling with the awkward and inadequate set of Roman numerals that hinder even the simplest addition or subtraction. The rules for combining numbers to deduce other ratios have been reduced to concise form and were taught to us in grammar school under the name of arithmetic.

Whereas the statements of the rules have become part and parcel of our mental processes, it is sometimes a painful operation to translate from John's apples or Mr. B's acres into gram-ionic weights or coulombs of electricity. Most of the computations of everyday life can be handled by the processes of mental arithmetic; for example, the cost of a week-end trip or the trade-in value of a radio which depreciates ten per cent annually. For scientific computations the simple methods will not always suffice. It will be well, therefore, to consider some of the aspects of what might be termed scientific arithmetic.

It should be recognized that while arithmetic is the simplest branch of mathematics, all applications of mathematics to scientific problems, however complex, lead up to arithmetic as the final operation. Regardless of the symbols, and of the reasoning processes by which the symbols are manipulated to produce a scientific formula, the usefulness of that formula is in computing results. Just as soon as specific numerical values replace the symbols in a formula, then the rest is arithmetic.

EXPONENTIAL NUMBERS

Most measurements in physical science result in either very large or very small values, and special methods of handling these numbers must be considered. The accepted value for the number of molecules in a gram-molecular weight of any substance is 606,000,000,000,000,000,000. To perform even the simplest multiplication or division with such a number would require more patience than skill. On the other hand, the weight of lead chromate dissolved by a liter of water at 20° C is 0.000043 gram, a value too small to handle conveniently. The familiar method of counting out units, tens, hundreds, thousands, etc., suggests that these very large and very small numbers could be stated as the product of two factors one of which expresses the significant digits and the other the magnitude. The simplest and the most generally used system takes advantage of the properties of exponential numbers. The number 220 may be analyzed into 2.2×100 ; but 100 is the square of 10, or 10^2 . Hence, we may express 220 as 2.2×10^2 . The first part (2.2) is called the **digit factor**, and, for convenience, it is always given with *one digit to the left of the decimal point*. The second part (10^2) is called the **exponential factor**, and it is always a

whole-number power of ten. Let us review the values of powers of ten compared with our usual numbers.

million.....	1,000,000. = 10^6	one-tenth.....	0.1 = 10^{-1}
hundred thousand.....	100,000. = 10^5	one-hundredth.....	0.01 = 10^{-2}
ten thousand.....	10,000. = 10^4	one-thousandth.....	0.001 = 10^{-3}
thousand.....	1,000. = 10^3	one-ten-thousandth.....	0.0001 = 10^{-4}
hundred.....	100. = 10^2	one-hundred-thousandth.....	0.00001 = 10^{-5}
ten.....	10. = 10^1	one-millionth.....	0.000001 = 10^{-6}
one.....	1. = 10^0		

The value 2,200,000 is 2.2 times one million, or 2.2×10^6 . It may be seen that the power of ten in the exponential factor is equal to the number of places the decimal point must be moved to derive the digit factor (remembering the definition of the digit factor). This rule also applies to the very small numbers. Thus, the solubility of lead chromate would be given as 4.3×10^{-5} gram per liter. Here the power of ten is negative. The complete rule is: the power of ten in the exponential factor is given by the number of places the decimal point must be moved to derive the digit factor counting *positive to the left* and *negative to the right*. Accordingly, the number of molecules in a gram-molecular weight is 6.06×10^{23} . Count and see for yourself.

By the use of exponential numbers the time and labor of counting figures to the right or left of the decimal point can be eliminated. Moreover, exponential numbers require less space and have an advantage in that the exponent gives an immediate insight into the magnitude of the number. Many of the data in chemical science (equilibrium constants, reaction velocities, solubilities, the sizes of molecules, etc.) are either very large or very small values, and these are seldom obtained with a precision extending to the fourth significant figure.

PROBLEM SET I

1. Change the following ordinary numbers to exponential numbers: (a) 67,500.; (b) 341.; (c) 9,652,170,000,000.; (d) 0.0000000007; (e) 0.003026; (f) 0.0201.

2. Convert these exponential numbers to ordinary numbers: (a) 4×10^{18} ; (b) 3.67×10^{-6} ; (c) 4.028×10^{-1} ; (d) 1.4×10^5 ; (e) 3.9×10^0 ; (f) 7.3145×10^8 .

3. Convert to exponential numbers: (a) 0.0000230; (b) 765.2; (c) 8.93; (d) 74,000,000.; (e) 0.0001001; (f) 0.05.

4. Convert to ordinary numbers: (a) 1.28×10^{-4} ; (b) 4.0×10^{-2} ; (c) 8.5031×10^7 ; (d) 2×10^{-10} ; (e) 3.12×10^{-2} ; (f) 1.03×10^0 .

Multiplication and Division of Exponential Numbers. In multiplication or division with exponential numbers, the digit factors are manipulated according to the rules of ordinary arithmetic. The exponential factors, however, must be treated according to the *law of algebraic summation of exponents*. This simply means that exponents are added taking cognizance of the signs. A few examples will illustrate.

$$(1) \quad 10^5 \times 10^2 = 10^{5+2} = 10^7$$

$$(2) \quad 10^{-4} \times 10^{-8} = 10^{-4+(-8)} = 10^{-12}$$

$$(3) \quad 10^5 \times 10^{-7} = 10^{5+(-7)} = 10^{-2}$$

It must also be remembered that a negative exponent indicates a reciprocal, or

$$10^{-2} = \frac{1}{10^2}$$

Conversely,

$$10^4 = \frac{1}{10^{-4}}$$

Taking now an example of multiplication of two exponential numbers:

$$2.0 \times 10^{-4} \times 1.8 \times 10^7$$

first collect digit factors and then exponential factors

$$2.0 \times 1.8 \times 10^{-4} \times 10^7$$

Next perform the operations separately

$$2.0 \times 1.8 = 3.6$$

$$10^{-4} \times 10^7 = 10^3$$

and combine the results into an exponential number, 3.6×10^3 .

In division, a similar routine may be followed:

$$\frac{4.8 \times 10^{-8}}{1.2 \times 10^3} = \frac{4.8}{1.2} \times 10^{-8} \times 10^{-3}$$

$$\frac{4.8}{1.2} = 4.0$$

$$10^{-8} \times 10^{-3} = 10^{-11}$$

and the answer is 4.0×10^{-11} .

Frequently the answer will not be in the form of a correctly expressed exponential number as we have defined its meaning. For example:

$$\frac{2.7 \times 10^{-12}}{3.0 \times 10^{-9}} = \frac{2.7}{3.0} \times 10^{-12} \times 10^9 = 0.90 \times 10^{-3}$$

Here the digit factor is not in the correct form, and the whole number must be translated. This operation must be carried out in detail to avoid error.¹ First determine what must be done to the digit factor to give it the correct form, and then perform the *reverse* operation on the exponential factor. In the above case the decimal point must be moved one place to the right to give a digit factor with one digit to the left of the decimal point. This necessitates multiplying the digit factor by 10.

$$0.90 \times 10 = 9.0$$

In order to restore the former value to the exponential number, the exponential factor must be divided by 10:

$$\frac{10^{-3}}{10} = 10^{-3} \times 10^{-1} = 10^{-4}$$

and the answer properly expressed is 9.0×10^{-4} . The net result of the two operations is to multiply the exponential number by $\frac{10}{10}$, which does not change its value. Two additional examples are:

$$(1) \quad 26.3 \times 10^7 = \frac{26.3}{10} \times 10^7 \times 10 = 2.63 \times 10^8$$

$$(2) \quad 0.071 \times 10^{14} = 0.071 \times 100 \times \frac{10^{14}}{100} = 7.1 \times 10^{12}$$

PROBLEM SET II

1. Solve:

$$(a) \quad 3.7 \times 10^{-3} \times 2.7 \times 10^{-4};$$

$$(c) \quad 7.0 \times 10^5 \times 1.2 \times 10^{11};$$

$$(e) \quad 1.5 \times 10^{-10} \times 4.3 \times 10^{10};$$

$$(b) \quad 4.1 \times 10^{16} \times 5.0 \times 10^{-3};$$

$$(d) \quad 1.25 \times 10^{-3} \times 3.10 \times 10^{-9};$$

$$(f) \quad 6.2 \times 10^{16} \times 1.3 \times 10^7.$$

2. Solve:

$$(a) \quad \frac{6.3 \times 10^{-17}}{2.1 \times 10^4};$$

$$(b) \quad \frac{7.2 \times 10^{19}}{1.2 \times 10^{16}};$$

$$(c) \quad \frac{1.21 \times 10^{-8}}{1.10 \times 10^2};$$

¹ The student is cautioned strongly against both short-cut methods and rules for the performance of this particular operation.

$$(d) \frac{8.1 \times 10^{-32}}{9.0 \times 10^{-37}}; \quad (e) \frac{9.9 \times 10^4}{3.3 \times 10^{-26}}; \quad (f) \frac{1.62 \times 10^{67} \times 6.80 \times 10^{-4}}{8.10 \times 10^{75}}.$$

3. Translate to correctly expressed exponential numbers:

$$(a) 162 \times 10^4; \quad (b) 216.2 \times 10^{-7}; \quad (c) 31.13 \times 10^{18}; \\ (d) 0.42 \times 10^{-8}; \quad (e) 0.0019 \times 10^6; \quad (f) 0.026 \times 10^{17}.$$

4. Solve:

$$(a) \frac{2.1 \times 10^{12} \times 7.0 \times 10^{-12}}{3.5 \times 10^{-7}}; \quad (b) \frac{16.4 \times 3.2 \times 10^{-9}}{4.0 \times 10^{-6} \times 0.32}; \\ (c) \frac{1200 \times 0.0075 \times 3.6 \times 10^{22}}{15 \times 6.0 \times 10^{19}}.$$

Powers and Roots of Exponential Numbers. The processes of raising to powers (evolution) or taking roots (involution) of exponential numbers represent merely an extension of the principles just considered. To raise an exponential number to a given power, the two factors are operated on separately. For instance:

$$(2.0 \times 10^5)^2 = (2.0)^2 \times (10^5)^2$$

The digit factor is squared in the usual way. The exponential factor may be written:

$$(10^5)^2 = 10^5 \times 10^5 = 10^{5+5} = 10^{10}$$

The result is simply a multiplication of the exponent by the power to which the factor is raised. The answer is:

$$4.0 \times 10^{10}$$

Likewise:

$$(3.0 \times 10^{-7})^3 = (3.0)^3 \times (10^{-7})^3$$

The digit factor is cubed in the usual fashion, the result, of course, being 27. The exponential factor may be written:

$$(10^{-7})^3 = 10^{-7} \times 10^{-7} \times 10^{-7} = 10^{(-7)+(-7)+(-7)} = 10^{-21}$$

The answer, translated into a correctly expressed exponential number, is 2.7×10^{-20} .

In taking roots of exponential numbers it is necessary to recall that taking the square root is the same as raising to the one-half power and taking the cube root is equivalent to raising to the one-

third power. Proceed then as in the previous case, performing the operations separately on the two factors.

$$\sqrt{9.0 \times 10^8} = (9.0 \times 10^8)^{\frac{1}{2}} = (9.0)^{\frac{1}{2}} \times (10^8)^{\frac{1}{2}} = 3.0 \times 10^4$$

And:

$$\begin{aligned}\sqrt[3]{8.0 \times 10^{-12}} &= (8.0 \times 10^{-12})^{\frac{1}{3}} \\ &= (8.0)^{\frac{1}{3}} \times (10^{-12})^{\frac{1}{3}} = 2.0 \times 10^{-4}\end{aligned}$$

In many instances, the product of the exponent and the fractional power is not a whole number, in which case the answer would be incomprehensible. It is only necessary to reverse the procedure for translating improperly expressed numbers to correctly expressed exponential numbers. Simply transform the number so that the exponent when multiplied by the fractional power gives a whole number (merely another way of saying that the exponent must be divisible by the number expressing the root). For example, $\sqrt[3]{2.7 \times 10^{13}}$ could be solved in either of two ways:

$$(2.7 \times 10^{13})^{\frac{1}{3}} = (27 \times 10^{12})^{\frac{1}{3}}$$

or

$$(2.7 \times 10^{13})^{\frac{1}{3}} = (0.027 \times 10^{15})^{\frac{1}{3}}$$

By the first transformation

$$(27)^{\frac{1}{3}} \times (10^{12})^{\frac{1}{3}} = 3.0 \times 10^4$$

and by the second transformation

$$\begin{aligned}(0.027)^{\frac{1}{3}} \times (10^{15})^{\frac{1}{3}} &= 0.30 \times 10^5 \\ &= 3.0 \times 10^4\end{aligned}$$

which is exactly the same answer as given in the first case.

Addition and Subtraction of Exponential Numbers. Cases of addition and subtraction of exponential numbers are less frequent, but it is necessary to consider them. It is obvious that $10^{27} - 10^{-14}$ would have no significance. The result would be comparable to the effect of removing one grain of sand from the shores of the Atlantic Ocean upon the remaining number of grains. The cases which are of significance are those in which addition or subtraction is performed upon exponential numbers of approximately the same magnitude, i.e., exponents approximately the same.

When the exponents are the same the operation is very simple. For example,

$$3.6 \times 10^{-4} - 2.1 \times 10^{-4}$$

may be factored and rearranged into

$$(3.6 - 2.1) \times 10^{-4} = 1.5 \times 10^{-4}$$

Similarly,

$$8.5 \times 10^5 + 1.3 \times 10^5 = (8.5 + 1.3) \times 10^5 = 9.8 \times 10^5$$

In case the exponents are not the same, the exponential numbers must be transformed so that all will be of the same magnitude.

For instance,

$$1.25 \times 10^6 + 2.50 \times 10^5$$

must be transformed to

$$1.25 \times 10^6 + 0.25 \times 10^6 = (1.25 + 0.25) \times 10^6 = 1.50 \times 10^6$$

This same transformation is required for subtraction.

PROBLEM SET III

1. Solve: (a) $(6.0 \times 10^{-4})^2$; (b) $(1.2 \times 10^7)^3$; (c) $(2.0 \times 10^{12})^4$; (d) $(0.0006)^3$.

2. Solve: (a) $\sqrt{6.4 \times 10^{-11}}$; (b) $\sqrt[3]{6.4 \times 10^{10}}$; (c) $\sqrt{8.1 \times 10^{13}}$; (d) $(3.6 \times 10^{-13})^{\frac{1}{2}}$.

3. Solve: (a) $1.5 \times 10^{-4} + 2.8 \times 10^{-4}$; (b) $4.2 \times 10^6 + 2.4 \times 10^6$; (c) $1.6 \times 10^{-7} + 3.2 \times 10^{-7} - 2.4 \times 10^{-7} - 1.3 \times 10^{-7}$; (d) $7.8 \times 10^{-3} + 2.0 \times 10^{-4} + 9.8 \times 10^{-4}$.

4. Solve: (a) $(1.7 \times 10^4 + 8.0 \times 10^3 - 2.0 \times 10^4)^2$; (b) $\sqrt{9.6 \times 10^{-13} + 2.0 \times 10^{-14} - 4.9 \times 10^{-13}}$; (c) $(3.0 \times 10^{-6} \times 5.0 \times 10^{13} - 1.5 \times 10^{-2} \times 4.0 \times 10^9)^{\frac{1}{2}}$.

LOGS (LOGARITHMS)

The complexity of the computations involved in astronomy and trigonometry as they had developed by the close of the sixteenth century had led mathematicians to search for simplifications of the arithmetical processes involving very large numbers. The most notable advance was made by John Napier (1550–1617), Laird of Merchiston, who announced the invention of logarithms² in his *Mirifici Logarithmorum Canonis Descriptio* printed in 1614. It may be remarked that Napier's "invention" was probably not

² *Encyclopædia Britannica*, 11th Edition, Cambridge University Press, 1911.

entirely original since the method seems to have been known to Euclid and Archimedes, but Napier invented the practical application and calculated the first tables of logarithms of trigonometric functions. Napier's³ announcement caused little response among the scientists of Europe. One of the first to recognize the importance and practical value of the new system was Henry Briggs (1556–1631). Briggs also recognized the inconvenience of Napier's "radix," which has a peculiar value,⁴ and therefore adopted what in modern notation amounts to the base 10. The tables of logarithms of numbers published by Briggs in 1617 were responsible for the immediate and widespread adoption of logarithms as a computation aid. The tables of Briggs and Vlacq, *Arithmetica Logarithmica* (1628), are substantially the same as those used today. Their work was so carefully done that no significant improvements have been effected. Logarithms have fittingly been termed the only great contribution to arithmetic in modern times.

In 1915, Chappel⁵ proposed the abridgment of the lengthy and awkward word "logarithm" to **log** with plural **logs**. His proposal bore no intention of introducing slang into the vocabulary of mathematics. Justification for the change arises from a consideration of the etymology of the name, which comes from the Greek words λόγος + αριθμός, meaning "the ratio of the numbers." Napier had demonstrated his principles by means of what he called the *ratio* between a geometric and an arithmetic series. The so-called *ratio* characterizes a relation between the series and not between the numbers, hence the full name does not carry the true import. Rather than attempt to adopt a new name, the abridgment seems more practicable. Hence the authors urge

³ Ball: *A Short History of Mathematics*, Macmillan, New York, 1893; and Cajori: *A History of Mathematics*, Second Edition, Macmillan, New York, 1919.

⁴ The original logarithms of Napier were related to the natural numbers they represented by the expression

$$N = 10^7 e^{-\log N / 10^7}$$

whereas the modern Napierian logarithm is given by

$$N = e^{\log N}$$

Though $e(2.71828)$ is the present base of Napierian logarithms, it is incorrect to state that Napier originally used it as "radix." His "radix" was actually

$$\left(1 - \frac{1}{10^7}\right)^{10^7}, \text{ which nearly equals } e^{-1}.$$

⁵ Chappel: *Five-Figure Mathematical Tables*, p. iv, W. and R. Chambers, Ltd., London, 1915.

the adoption of Chappel's suggestion, and throughout the remainder of this book the terms **log** and **logs** will be used.

Computations with logs represent merely a further extension of the foregoing law of summation of exponents. Recalling that:

$$10 = 10^1$$

$$3.162 = \sqrt{10} = 10^{0.5}$$

$$2.154 = \sqrt[3]{10} = 10^{0.33}$$

$$1 = 10^0$$

it can readily be seen that numbers between 10 and 1 can be expressed as exponential factors in which the exponents have values between 1 and 0. Given a means of determining what fractional exponent of 10 gives an exponential number equivalent to a digit factor, the process of multiplication of digit factors becomes simply an addition of the fractional exponents. Likewise, division of one digit factor by another becomes mere subtraction of exponents. We have already seen that multiplication and division of ordinary exponential factors involve only addition and subtraction, respectively, of the exponents. Further, since any number may be expressed as an exponential number of two factors, all multiplication may be reduced to adding exponents and all division to subtracting exponents. Values of fractional exponents corresponding to values of digit factors are given in Table I in the Appendix. Let us first express a number as an exponential factor. From the table,

$$7.30 = 10^{.8633}$$

If we take an exponential number such as 7.30×10^5 , this could be written

$$7.30 \times 10^5 = 10^{.8633} \times 10^5$$

Adding the two exponents together we get a composite exponential factor which represents both the factors of the exponential number

$$7.30 \times 10^5 = 10^{5.8633}$$

This composite exponent is the log of the exponential number 7.30×10^5 . The log consists of two parts: the **mantissa** (.8633), which is the fractional exponent corresponding to the digit factor; and the **characteristic** (5.), which expresses the magnitude. The following table will impress these definitions.

Number	Exponential Number	Composite Exponential Factor	Log
7.30	7.30×10^0	$10^{0.8633}$	0.8633
73.0	7.30×10^1	$10^{1.8633}$	1.8633
730.	7.30×10^2	$10^{2.8633}$	2.8633
7,300.	7.30×10^3	$10^{3.8633}$	3.8633
73,000.	7.30×10^4	$10^{4.8633}$	4.8633
730,000.	7.30×10^5	$10^{5.8633}$	5.8633

When a number less than 1 is expressed as an exponential number, the exponent, of course, has a negative value. In this instance, $0.0730 = 7.30 \times 10^{-2}$, which may be rewritten $10^{.8633} \times 10^{-2}$. To express this as a composite exponential factor, in order to preserve the significance of the characteristic as negative and of the mantissa as positive, we write the negative sign above the 2 as a superbar, thus $\bar{2}$. This convention allows us then to write $\bar{2}.8633$ as the log of 0.0730. But it must be kept constantly in mind that this value $\bar{2}.8633$ is an artificial numerical expression and not the same thing as either -2.8633 or $+2.8633$.

The manipulation of logs as presented in elementary mathematics courses involves the statement of a log as the difference of two terms. Giving the log of 0.0730 as $8.8633 - 10$ shows a positive mantissa and a characteristic which equals $8 - 10$, or -2 . Therefore, no new principle is being introduced, only a different notation. It has been our experience that the traditional presentation of logs leads to greater possibility of error in chemical computations than the system being presented here.

Let us examine the following table, which is similar to the one at the top of this page.

Number	Exponential Number	Composite Exponential Factor	Log
0.730	7.30×10^{-1}	$10^{\bar{1}.8633}$	$\bar{1}.8633$
0.0730	7.30×10^{-2}	$10^{\bar{2}.8633}$	$\bar{2}.8633$
0.00730	7.30×10^{-3}	$10^{\bar{3}.8633}$	$\bar{3}.8633$
0.000730	7.30×10^{-4}	$10^{\bar{4}.8633}$	$\bar{4}.8633$

Multiplication by Use of Logs. A few examples of multiplication by the use of logs will demonstrate the process.

$$\begin{aligned}\text{Example 1. } 652. \times 127. &= 6.52 \times 10^2 \times 1.27 \times 10^2 \\ 6.52 \times 10^2 &= 10^{.8142} \times 10^2 = 10^{2.8142} \\ 1.27 \times 10^2 &= 10^{.1038} \times 10^2 = 10^{2.1038}\end{aligned}$$

The operation then becomes:

$$10^{2.8142} \times 10^{2.1038}$$

The sum of the logs is: $\log 652 = 2.8142$

$$\log 127 = 2.1038$$

$$\log (652 \times 127) = 4.9180$$

In order to interpret the sum of the logs we can write:

$$652 \times 127 = 10^{4.9180} = 10^{.9180} \times 10^4$$

From the table of logs we find that the fractional exponent corresponds to a digit factor (reversing the previous process of reading the table) equal to 8.28. The answer then is 8.28×10^4 , or 82,800. Actual multiplication gives 82,804.

$$\begin{aligned}\text{Example 2. } 3.21 \times 0.00295 &= 3.21 \times 10^0 \times 2.95 \times 10^{-3} \\ 3.21 \times 10^0 &= 10^{0.5065} \\ 2.95 \times 10^{-3} &= 10^{\bar{3}.4698}\end{aligned}$$

The sum of the logs is

$$0.5065$$

$$\bar{3}.4698$$

$$\bar{3}.9763$$

Then,

$$10^{\bar{3}.9763} = 10^{.9763} \times 10^{-3} = 9.47 \times 10^{-3}$$

or

0.00947. Actual multiplication gives 0.0094695.

$$\begin{aligned}\text{Example 3. } 53.9 \times 690. \times 0.753 &= 5.39 \times 10^1 \times 6.90 \times 10^2 \times 7.53 \\ \times 10^{-1} &= 10^{.7316} \times 10^1 \times 10^{.8388} \times 10^2 \times 10^{.8768} \times 10^{-1} = 10^{1.7316} \\ \times 10^{2.8388} \times 10^{\bar{1}.8768}\end{aligned}$$

The sum of the logs is

$$1.7316$$

$$2.8388$$

$$\bar{1}.8768$$

$$4.4472$$

The answer is $10^{4.4472} = 10^{.4472} \times 10^4 = 2.80 \times 10^4 = 28,000$. Actual multiplication gives 28,005.

In the three preceding examples, comparison of the answers by logs with the answers by longhand computation shows slight but

still noticeable differences. It is now time to point out that the values for the mantissas given in tables of logs are only approximate. In the language of mathematics, the mantissas are irrational numbers. The complete expression of such an irrational number would require an infinite number of figures to the right of the decimal point. From the practical standpoint it is obvious that a limit to the number of figures must be set, and experience has shown that four or five figures in the mantissa are sufficient for the general purposes of the chemist. The great advantage in computing by logs is that, after some experience has been obtained, the steps become almost automatic and much time is saved.

Consider these slightly more complicated examples of multiplication.

$$\text{Example 4. } 392 \times 1.67 \times 72.8 \times 0.00294 = 3.92 \times 10^2 \times 1.67 \times 10^0 \times 7.28 \times 10^1 \times 2.94 \times 10^{-3}$$

$$\begin{aligned}\log 392 &= 2.5933 \\ \log 1.67 &= 0.2227 \\ \log 72.8 &= 1.8621 \\ \log 0.00294 &= \bar{3}.4683\end{aligned}$$

$$2.1464$$

Referring to the log table it is seen that the mantissa corresponds to a value between 1.40 and 1.41. The difference between the values .1461 and .1492 is .0031. Our mantissa falls $\frac{.0003}{.0031} = \frac{3}{31}$ of the way between 1.40 and 1.41. The fraction $\frac{3}{31}$ is approximately $\frac{1}{10}$; therefore the correct answer is 1.401×10^2 . This method by which intermediate values are found is called *interpolation*. A full discussion of interpolation will be found in Chapter V.

$$\text{Example 5. } 4.16 \times 9.81 \times 0.528 \times 86.2 = 4.16 \times 10^0 \times 9.81 \times 10^0 \times 5.28 \times 10^{-1} \times 8.62 \times 10^1$$

$$\begin{aligned}\log 4.16 &= 0.6191 \\ \log 9.81 &= 0.9917 \\ \log 0.528 &= \bar{1}.7226 \\ \log 86.2 &= 1.9355\end{aligned}$$

$$3.2689$$

Again it is found that the mantissa falls between two values in the table. The difference between the values .2672 and .2695 is .0023. Our value is $\frac{.2689 - .2672}{.0023} = \frac{.0017}{.0023} = \frac{17}{23}$ of the way between 1.85 and 1.86. The fraction $\frac{17}{23}$ approximately equals $\frac{7}{10}$. The correct answer is therefore $1.857 \times 10^3 = 1,857$.

Division by Use of Logs. In division with logs, the artificial numerical character of logs must be kept in mind while subtracting. The law upon which all this is based, it will be remembered, is the law of algebraic summation of exponents. Regard for the sign must of necessity always be foremost. A few examples carefully followed through will illustrate.

$$\text{Example 1. } \frac{145}{25.0} = \frac{1.45 \times 10^2}{2.50 \times 10^1} = \frac{10^{2.1614}}{10^{1.3979}}$$

$$\begin{array}{r} \log 145 = 2.1614 \\ -\log 25.0 = -1.3979 \\ \hline 0.7635 \end{array}$$

Since in this case both the characteristic and mantissa are positive, placing the negative sign before the log makes both parts negative. Referring to the log table, our mantissa is found to fall between two values. The process of interpolation gives: $\frac{.7635 - .7634}{.7642 - .7634} = \frac{.0001}{.0008} = \frac{1}{8}$. The fraction is nearer $\frac{1}{10}$ than $\frac{2}{10}$. The answer is therefore $5.801 \times 10^0 = 5.801$.

$$\text{Example 2. } \frac{0.00136}{38.9} = \frac{1.36 \times 10^{-3}}{3.89 \times 10^1} = \frac{10^{\bar{3}.1335}}{10^{1.5899}}$$

$$\begin{array}{r} \log 0.00136 = \bar{3}.1335 \\ -\log 38.9 = -1.5899 \\ \hline \bar{5}.5436 \end{array}$$

The answer is 3.497×10^{-5} .

$$\text{Example 3. } \frac{69.5}{0.0528} = \frac{6.95 \times 10^1}{5.28 \times 10^{-2}}$$

$$\begin{array}{r} \log 69.5 = 1.8420 \\ -\log 0.0528 = -\bar{2}.7226 \\ \hline 3.1194 \end{array}$$

The answer is 1,316. In this case the second log has a negative characteristic and a positive mantissa. Placing the negative sign before the log reverses the signs of the two parts, and the result is a positive characteristic and a negative mantissa. Algebraic summation gives the result shown.

Example 4. Choosing some logs at random, the results of subtraction are:

2.8639	4.3874	0.0531	14.9542
-1.5378	-0.9047	-2.9557	-6.5707
3.3261	5.4827	3.0974	20.3835

Operations Involving Both Multiplication and Division. In an operation involving both multiplication and division, it is best to sum first the logs of all factors in the numerator and then sum the logs of all factors in the denominator. The division is then performed by subtracting the second log from the first, just as in the previous section on division. Examples are:

$$\text{Example 1. } \frac{23.1 \times 0.317}{0.0116 \times 0.472} = \frac{2.31 \times 10^1 \times 3.17 \times 10^{-1}}{1.16 \times 10^{-2} \times 4.72 \times 10^{-1}}$$

log 23.1 = 1.3636	log 0.0116 = $\bar{2}.0645$
log 0.317 = $\bar{1}.5011$	log 0.472 = $\bar{1}.6739$
log numerator = 0.8647	log denominator = $\bar{3}.7384$
log numerator = 0.8647	
-log denominator = $-\bar{3}.7384$	
log answer = 3.1263	

The answer is 1.338×10^3 .

Example 2. Stripping a computation by logs to its bare essentials:

$\frac{312 \times 74.8}{0.0189 \times 6.73 \times 9.26}$	2.4942	$\bar{2}.2765$	4.3681
	1.8739	0.8280	-0.0711
	4.3681	0.9666	4.2970
		0.0711	

The answer is 1.981×10^4 .

Raising a Number to a Power by Use of Logs. Raising a number to a power by logs requires only three steps: (1) find the log of the number; (2) multiply the log by the power; (3) find the answer corresponding to the resultant log. Just as before, remember that the mantissa is always positive and the characteristic may be either positive or negative, and carry out the algebraic product accordingly.

$$\text{Example 1. } (21.7)^3 = (2.17 \times 10^1)^3 = (10^{1.3365})^3$$

1.3365
3
4.0095

The answer is 1.022×10^4 .

Example 2. $(0.0205)^2 = (2.05 \times 10^{-2})^2 = (10^{\bar{2}.3118})^2$

$$\begin{array}{r} \bar{2}.3118 \\ 2 \\ \hline \bar{4}.6236 \end{array}$$

The answer is 4.204×10^{-4} .

Example 3. $(0.341)^3 = (3.41 \times 10^{-1})^3 = (10^{\bar{1}.5328})^3$

$$\begin{array}{r} \bar{1}.5328 \\ 3 \\ \hline \bar{2}.5984 \end{array}$$

The answer is 3.967×10^{-2} .

Extracting the Root of a Number by Use of Logs. Extracting the root of a number when the characteristic is positive presents no new problem. When the root is expressed as a fractional power, the log of the number is multiplied by the fraction, or, to say it another way, the log is divided by the integer expressing the root.

Example 1. $\sqrt{62.5} = (6.25 \times 10^1)^{1/2} = (10^{1.7959})^{1/2} = 10^{1.7959/2}$
 $= 10^{0.8980}$

The answer is 7.907.

Example 2. $\sqrt[3]{4.37} = (10^{0.6405})^{1/3} = 10^{0.6405/3} = 10^{0.2135}$

The answer is 1.598.

Extraction of the root of a number when the characteristic is negative does introduce new aspects of the law of algebraic summation of exponents. There is one simple case which we shall consider first. If the negative characteristic happens to be a simple multiple of the integer representing the root desired, then the operation is much like that in Examples 1 and 2.

Example 3. $\sqrt[3]{0.00000431} = (4.31 \times 10^{-6})^{1/3} = (10^{\bar{6}.6345})^{1/3}$
 $= 10^{\bar{6}.6345/3} = 10^{\bar{2}.2115}$

The answer is 1.627×10^{-2} .

Negative Logs. When the negative characteristic is not a simple multiple of the integer expressing the root, we must convert the log from its artificial character to a pure negative number (or a

negative log). The transformation is accomplished by a simple algebraic summation of the two parts of the log thus:

$$\text{Example 1. } \bar{3}.4000 = -3.0000 + 0.4000 = -2.6000.$$

$$\text{Example 2. } \bar{4}.6500 = -4.0000 + 0.6500 = -3.3500.$$

$$\text{Example 3. } \bar{2}.1826 = -2.0000 + 0.1826 = -1.8174.$$

A negative log so obtained may be divided by the integer representing the root. The result is a negative log. We must now reverse the transformation process to express the log in its artificial form as originally defined. This step is perhaps the most difficult one in the mastery of logs, and it must always be performed with great caution. The procedure is to add 1 to the mantissa and subtract 1 from the characteristic of the negative log. Such manipulation does not alter the value. The following examples will illustrate.

$$\text{Example 4. } -4.7000 = -4 - 1 + 1.0000 - 0.7000 = (-4 - 1) + (1.0000 - 0.7000) = \bar{5}.3000.$$

$$\text{Example 5. } -2.4600 = (-2 - 1) + (1.0000 - 0.4600) = \bar{3}.5400.$$

$$\text{Example 6. } -7.2372 = (-7 - 1) + (1.0000 - 0.2372) = \bar{8}.7628.$$

Now to perform the complete operation of extracting the root of a number, the log of which has a negative characteristic:

$$\begin{aligned} \text{Example 7. } \sqrt[3]{0.0438} &= (4.38 \times 10^{-2})^{1/3} = (10^{\bar{2}.6415})^{1/3} \\ &= 10^{-1.3685/3} = 10^{-0.4528} = 10^{\bar{1}.5472} \end{aligned}$$

The answer is 3.525×10^{-1} .

$$\begin{aligned} \text{Example 8. } \sqrt{0.00819} &= (10^{\bar{3}.9133})^{1/2} = 10^{-2.0867/2} = 10^{-1.0434} \\ &= 10^{\bar{2}.9566}. \end{aligned}$$

The answer is $9.06 \times 10^{-2} = 0.0906$.

The necessity for converting to negative logs is not by any means confined to the case just discussed. In many of the formulas of physical chemistry, the log of some factor must be multiplied by other numbers. It is readily apparent that a log with negative characteristic could not be set in the formula in its artificial form. The result would be meaningless. Conversion of the log to a negative log is arithmetically imperative.

Example 9. Almost any standard textbook of elementary physical chemistry will give the following formula for the calculation of the minimum applied emf required to begin electrodeposition of copper from a solution of a copper salt of concentration $C_{\text{Cu}^{++}}$ at 25°C :

$$E = 0.344 + \frac{0.0591}{2} \log C_{\text{Cu}^{++}}$$

For a solution containing 0.225 gram ion per liter of Cu^{++} ion, the minimum emf would be:

$$E = 0.344 + \frac{0.0591}{2} \log 0.225$$

$$\log 0.225 = \bar{1}.352 = -0.648$$

$$E = 0.344 + \frac{0.0591}{2} (-0.648)$$

$$= 0.344 - 0.0192 = 0.325 \text{ volt}$$

Example 10. The $p\text{H}$ of a solution, which gives directly a measure of the H^+ ion concentration and indirectly of the OH^- ion concentration (since $14 - p\text{H} = p\text{OH}$), is defined as the log of the reciprocal of the H^+ ion concentration expressed in gram ions per liter. That is,

$$p\text{H} = \log \frac{1}{C_{\text{H}^+}}$$

To find the value of C_{H^+} when the $p\text{H}$ is given involves the use of negative logs. For example, suppose the $p\text{H}$ of a solution to be 9.35. What is the H^+ ion concentration?

$$9.35 = \log \frac{1}{C_{\text{H}^+}} = \log 1 - \log C_{\text{H}^+}$$

But

$$\log 1 = 0$$

Therefore,

$$9.35 = -\log C_{\text{H}^+}$$

or

$$\log C_{\text{H}^+} = -9.35$$

Transforming the negative log to a log expressed in the usual artificial fashion:

$$(-9 - 1) + (1.00 - 0.35) = \bar{10}.65$$

and

$$C_{\text{H}^+} = 4.5 \times 10^{-10}$$

The student pursuing advanced courses in chemistry will find the use of logs a necessity. It may be pointed out that log tables giving mantissas to five figures permit greater accuracy in computation than is possible with four-place tables, and simplify greatly

the process of interpolation. Furthermore, it must be noted that log tables are rarely given with all the decimal points; as is the case in Table I (Appendix). This omission is merely a typographical convenience; it is universally understood that values of N are to be interpreted as properly expressed digit factors, and the mantissas are understood to be decimal fractions. Realizing this, one should find no difficulty in reading any log table. The student interested in pursuing the subject of logs further is advised to purchase a book of five-place log tables. He will find, in addition to trigonometric tables, tables of cologs (logs of reciprocals of numbers) which make division by a number a multiplication with the reciprocal of the number and eliminate subtraction of logs, and tables of antilogs (tables of numbers corresponding to values of mantissas) which simplify the translation of the resultant log to answers expressed as digit factors.⁶

Natural (Napierian) Logs. The integration of many of the differential equations of theoretical chemistry results in terms involving natural or Napierian logs. One such integrated term which often appears is

$$RT \ln K$$

The symbol $\ln$ has received widespread adoption among chemists to represent natural log. The mathematicians use

$$\log_e K$$

It is possible to derive a log system with any base; the general symbol

$$\log_a X$$

means the log to base a of X . Since in chemical computations only two bases, e and 10, are ever used, we can state the equivalence of our symbols to those of the mathematicians as follows:

Chemical Mathematical

$$\log X = \log_{10} X$$

$$\ln X = \log_e X$$

⁶ See footnote 5 on page 9.

To eliminate the necessity for two separate log tables, we make use of the relations

$$\ln X = 2.3026 \log X$$

and

$$\log X = 0.43429 \ln X$$

Whenever an integrated expression contains a natural log we make the substitution

$$RT \ln K = 2.3026 RT \log K$$

to permit computation with common logs. On the other hand, if an expression containing a Briggsian (or common) log is to be differentiated, we make the substitution

$$m \log X = 0.43429 m \ln X$$

An example of this latter substitution will be observed on page 61.

PROBLEM SET IV

1. It is suggested that the student repeat 1, 2, and 4 in Problem Set II and 1, 2, and 4 in Problem Set III, using logs.

2. Express the following H^+ ion concentrations as pH: (a) 3.2×10^{-4} ; (b) 7.9×10^{-12} ; (c) 4.32×10^{-6} .

3. Convert the following values of pH to OH^- ion concentrations: (a) 13.6; (b) 9.4; (c) 6.3.

4. The symbol p as in pH has received acceptance as an *operator* (see p. 69), signifying the log of the reciprocal of whatever function it prefixes.

(a) $K_{s.p.}$ of $Fe(OH)_3 = 1.1 \times 10^{-36}$. Calculate $pK_{s.p.}$

(b) $K_{(instab.)}$ of $Ag(NH_3)_2^+ = 6.8 \times 10^{-8}$. Calculate $pK_{(instab.)}$.

(c) $K_{(ion)}$ of $HCN = 7 \times 10^{-10}$. Calculate $pK_{(ion)}$.

5. Solve:

$$(a) 0.816 + \frac{0.0591}{2} \log \frac{0.250}{1.90}.$$

$$(b) \text{ Find } C_A \text{ when } 0.0529 = 0.0521 \log \frac{0.315}{C_A}.$$

THE SLIDE RULE

The slide rule is frequently called "a mechanical log table." This conception is amply justified by a study of the history of its conception and invention. Historians⁷ of mathematics now seem agreed that the first slide rule was constructed by William Oughtred

⁷ Cajori: *A History of the Slide Rule*, Engineering News Publishing Co., New York, 1909.

(1574–1660) by putting together two Gunter's rules, and that the account of his invention was first published in 1632. The rule of E. Gunter (1581–1626) was a mechanically inscribed Briggsian log table on which multiplications and divisions were performed with dividers. Oughtred also invented the circular slide rule. Various mechanical improvements on Oughtred's device resulted, but the only notable advances were made by J. Robertson who in 1778 constructed a slide rule with a runner, by P. M. Roget who in 1815 invented the log-log scale, and by A. Mannheim (1831–1906) who about 1850 devised the slide rule that with only slight modifications is in use today.

To reproduce Oughtred's invention crudely, and at the same time perform an enlightening and instructive experiment, each student should execute the following exercise:

- (a) Procure a sheet of graph paper, and cut from it a strip 10 divisions (cm) long and 2 divisions (cm) wide.
- (b) Lengthwise draw a line along the middle.
- (c) Along the top and along the bottom from left to right number the divisions (cm) from 0 to 1 at intervals of a tenth.
- (d) From the table of logs given in the Appendix obtain the values of the logs corresponding to the integers from 1 to 10.
- (e) Make a mark across the center line at the location of the log corresponding to each of the integers, and at each end of the mark write the integer. The strip should now look like Fig. 1.
- (f) Cut the strip along the center line.
- (g) To distinguish the strips and simultaneously conform to slide-rule convention, label the upper strip C and the lower D.

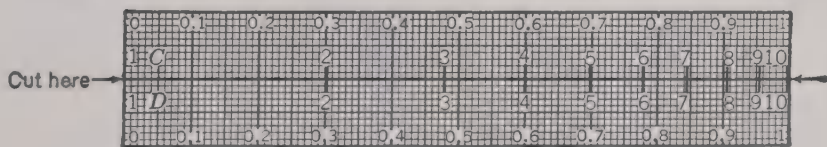


FIG. 1:

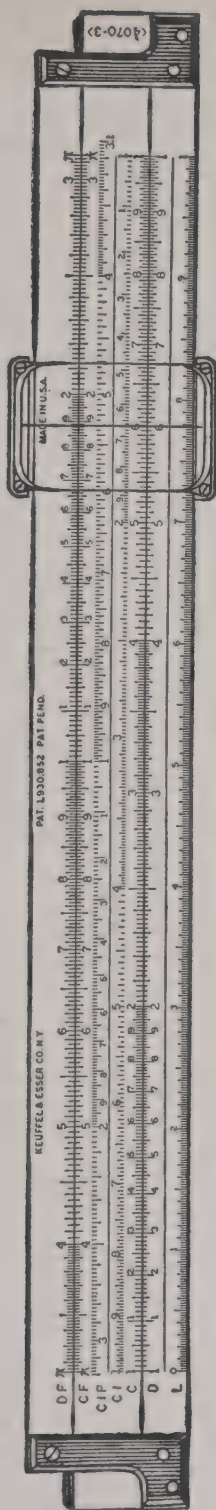
Let us now multiply 2 by 3. To do this move C along D until 1 on the C scale (called the left index) coincides with 2 on the D scale. Then look along the C scale to 3, and note that it coincides with 6 on the D scale. The net effect is to add mechanically the scale length of the log of 2 to the scale length of the log of 3. Since

the digit factor corresponding to the scale length represented by the sum of the logs is on the D scale the product of the numbers is read directly. Now try 3×3 ; 2×4 ; 2×5 . By the same steps multiply 4 by 5. The answer in this case lies beyond the right end of the D scale, and the interpretation is that the sum of the logs of 4 and 5 is greater than 1. To surmount this difficulty, curl strip D into a cylinder with the right and left indexes just touching, and grasp between the thumb and index finger of the left hand. Treat strip C in the same way, and hold with the right hand. Now, just as before, move C above D until the index of C is above 4 of D. Look along C to 5, and it will be found to coincide with 2 on the D scale. Obviously the answer must be 20. This is consistent with computations by logs because only the mantissas are being taken into account and as will be subsequently seen the characteristic must be cared for in another way. Now flatten out the two strips, and perform the same multiplication in this way. Set 10 on the C scale (called the right index) above 4 on D; then under 5 on C will be found 2 on D in perfect agreement with the previous case. The effect then of reversing the direction of the slide is to double the scale around on itself like a snake swallowing its tail, thereby producing a continuation of the scale.

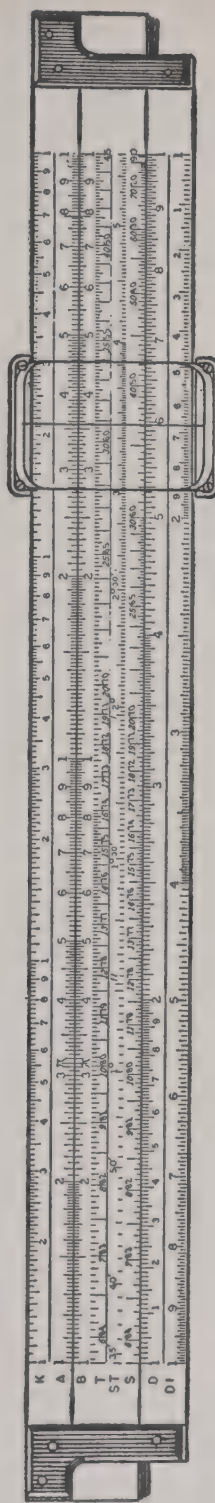
Division is simply a complete reversal of the process of multiplication. To divide 6 by 3, above 6 on D set 3 on C, and under the left index of C will be found on D the answer 2. This is merely a subtraction of the scale length of the log of 3 from the scale length of the log of 6. Try $8 \div 2$. Now try $4 \div 5$. The left index of C is off the scale, but the right index is above 8 on the D scale. The answer must be 0.8, and, as in a previous instance, this may be interpreted as a continuation of the scale by the same process of doubling back on itself.

Now that this little scene from mathematical history has been re-enacted, let us examine a modern slide rule of the "Duplex" type⁸ (Fig. 2). Except for more graduations, more scales, and

⁸ A slide rule with the linear log scale (L), in addition to the K, A, B, C, and D scales, is by far the most useful type for the chemist or chemical engineer. Each student is strongly urged to purchase a slide rule. Not only will a slide rule be helpful in subsequent exercises in this book, but also it will prove an efficient aid to computations throughout your professional career. The authors do not recommend the "chemical" slide rule to students of chemistry.



Front



Back

FIG. 2: Slide Rule.

Courtesy of Keuffel & Esser Co.

obvious mechanical improvements, it does not differ from our "play" slide rule. The C and D scales appear in the same positions. Concentrating for the time on only the C and D scales, repeat all the multiplications and divisions previously performed with the two paper strips. Note next the spacing of the graduations between the integers. Between 1 and 2 the smallest division represents 0.01; between 2 and 4 the smallest division represents 0.02; and between 4 and 10 the smallest division represents 0.05. With the valuations of the graduations in mind, perform the following operations, writing down the answers: (1) 1.3×4.5 ; (2) 2.6×2.96 ; (3) 1.6×3.5 ; (4) 3.7×6.0 ; (5) $3.62 \div 1.81$; (6) $4.7 \div 2.7$. Compare your answers with the correct ones: (1) 5.85; (2) 7.7; (3) 5.5; (4) 2.22; (5) 2.0; (6) 1.74.

For more complicated examples, use is made of the runner with engraved hairline. By setting the hairline on a figure the necessity of writing it down or attempting to remember it is obviated. To solve

$$\frac{3.67 \times 2.41}{4.18}$$

first perform the multiplication as usual by setting the left index of C above 3.67 on D, and the product 8.85 is seen on D below 2.41 on C. Set the hairline of the runner on the product 8.85, and perform the division by setting 4.18 on the C scale to coincide with the hairline. The answer 2.12 is found on D under the left index of C. To solve

$$\frac{1.65 \times 2.79 \times 2.12}{4.38 \times 1.03}$$

perform the multiplication of the three factors in the numerator, using the runner, and divide successively by each of the two factors of the denominator, making use of the runner also to mark the dividend after the first division. The result is 2.16.

It has already been indicated that the slide rule can handle only the digit factors. The characteristics must be cared for by the familiar method of expressing the numbers as exponential numbers rounded off to give an approximate answer. To solve

$$\frac{862 \times 4.73 \times 0.0118}{942 \times 0.246}$$

rewrite it thus:

$$\frac{9 \times 10^2 \times 5 \times 10^0 \times 1 \times 10^{-2}}{9 \times 10^2 \times 3 \times 10^{-1}} = \frac{9 \times 5 \times 1}{9 \times 3} \times 10^{-1} \\ = \frac{45}{27} \times 10^{-1} = 0.2 \text{ (approx.)}$$

The slide rule gives 208, and by the rough calculation it is apparent that the decimal must be placed to give 0.208 as the answer.

Powers and Roots. Squaring and cubing numbers can be accomplished on the slide rule with the use of only the runner, leaving the slide in its normal position. Compare the A scale with D. It is apparent that the A scale is made up of two complete scales from 1 to 10 each identical to the other. The length from 1 to 10 on the A scale is the same as from 1 to 3.162 ($\sqrt{10}$) on the D scale and the same as from 0 to 0.5 on the linear log scale L. The log scale corresponding to scale A (only the log scale of D is ever shown) must therefore be equivalent to twice the L scale corresponding to D. If the hairline of the runner is placed on a number on the D scale, the number given by the hairline on A must be the square of the number on D ($\log N^2 = 2 \log N$).

Similarly the K scale is seen to correspond to an unshown log scale that is three times the log scale L. A setting of the runner on the D scale gives the cube of the number by reading the K scale.

Square roots are obtained by reversing the above procedure. Set the runner on the number on the A scale, and its square root is found on D. A rough computation according to page 7, taking account of odd and even exponents, will show on which half of the A scale to set the runner. For instance:

$$\sqrt{4.0 \times 10^3} = \sqrt{40 \times 10^2} = \sqrt{40} \times 10$$

If the runner were set on the first 4 on A, the hairline would show 2 on D. The square root of 40 is between 6 and 7; therefore the second 4 must be the setting on A. The answer is $6.33 \times 10 = 63.3$. This treatment can be extended to the problem of magnitude in taking cube roots.

The cube root is obtained by setting the runner on the number on K and reading the cube root on D.

Logs. In solving expressions resulting from substituting numerical values in the Nernst equation for the emf of a galvanic cell, such terms arise as:

$$\frac{0.0591}{2} \log 0.225$$

It is best to find first the value of the log term by setting the runner at 2.25 on the D scale and reading the mantissa on the L scale. The mantissa is 0.352, and the complete log is $\bar{1}.352$. Expressed as a negative log, this is -0.648 . Now the term becomes

$$\frac{0.0591 (-0.648)}{2} = -\frac{5.91 \times 6.48}{2} \times 10^{-3} = -0.0192$$

Compare with Example 9 on page 18.

The following exercises will give additional practice in slide-rule manipulation, but it must be remembered that only practice (and a great deal of it) can make perfect. Many operations, other than the ones described, can be performed on the various types of slide rules. Their use is so infrequently required by the chemist that no attempt will be made here to describe them. Manuals supplied by the manufacturers describe the more complex operations in a manner that will be clear to one who has already mastered the more elementary operations.

PROBLEM SET V

1. Solve the following with the slide rule: (a) $42.6 \div 0.204$; (b) 14.7×6.38 ;

(c) $\frac{138. \times (4.13)^2}{22.6}$; (d) $\sqrt{\frac{8.23}{9.81}}$; (e) $72.9 \sqrt[3]{9.42}$.

2. Compute with the slide rule values of y for the given values of x from the equation:

$$y = \left(\frac{42.3x}{128} \right)^2$$

(a) $x = 0.879$; (b) $x = 3.67$; (c) $x = 8.25$; (d) $x = 15.9$.

3. Solve with the slide rule: (a) $\sqrt{26.9}$; (b) $\sqrt[3]{9.19}$; (c) $\sqrt{4.36 \times 1.93}$;

(d) $\sqrt{\frac{42.1}{8.34} \log 6.72}$.

4. Solve with the slide rule: (a) $\frac{71.4 \times 478}{89.3 \times 32.8}$; (b) $\frac{(3.29)^3 \times 14.9}{8.20 \log 4.18}$;

(c) $\frac{1.85 \times 10^{-5} \times 0.525}{0.817}$; (d) $\frac{248 \sqrt{37.1}}{16.7 \times 22.2}$.

SUPPLEMENTARY PROBLEMS

1. The average coefficient of expansion of gases is $1/273$ per degree Centigrade. Express this value as an exponential number.

2. The "average speed" of the hydrogen molecule, under standard conditions, is 184,000 cm per sec; that of the oxygen molecule is 46,000 cm per sec. Express these speeds exponentially.

3. The volume of the hydrogen molecule has been calculated to be 4.8×10^{-23} cubic centimeter. Express the volume as a decimal fraction.

4. The number of molecules of gas per liter at S.T.P. has been determined by several reliable methods and found to be 2.70×10^{22} , with a probable error of less than one per cent. Write this value as an ordinary number.

5. Since the number of molecules in 1 cc of a gas is 2.70×10^{19} , by taking the cube root of this we obtain the number in one linear centimeter. Make this computation. What is the *average spacing of molecules*, i.e., the average distance apart from center to center? The *actual* distance apart of any two molecules varies and is constantly changing.

6. The *mean free path in gases* is the average distance molecules travel between collisions. This must not be confused with the average spacing of molecules (calculated in Problem 5). The mean free path for any gas under S.T.P. was first calculated by James Clerk Maxwell (1860) and is not far from 0.00001 cm. Express this in the exponential form. How many of its neighbors will a molecule pass before finally colliding with one? *Hint:* The average spacing of molecules was computed in Problem 5.

7. At room temperature, a hydrogen molecule collides with other hydrogen molecules about 10,000,000,000 times per second! Write this value as an exponential number.

8. Millikan's value for the electronic charge, e , is 4.774×10^{-10} electrostatic unit. Express this as a decimal fraction. Several years ago Millikan's value of e was corrected, upon the basis of new measurements of the viscosity of air. The revised value is 4.78×10^{-10} . Express the correction as an exponential number.

9. Planck's universal constant, h , has a value of 6.55×10^{-27} erg-second. How many zeros would be required to write this number as a decimal fraction?

10. A micron (μ) is 0.001 mm, and a millimicron ($m\mu$) is 0.000001 mm. Express these units exponentially.

11. The mean wave length of the two yellow lines of sodium is 5893×10^{-7} mm. Express this value in Ångström units. The Ångström unit (Å) is 10^{-8} cm.

12. The ionization constant, K_i , of acetic acid is 0.000018 at room temperature. Express K_i exponentially.

13. The solubility-product constant, $K_{s.p.}$, of silver iodide is 1.5×10^{-16} . How many zeros are required to write the $K_{s.p.}$ as a decimal?

14. The instability constant (complex dissociation constant), $K_{(instab.)}$, of the silver ammonia complex ion, $Ag(NH_3)_2^+$, is 0.000000068. Write the $K_{(instab.)}$ value as an exponential number.

15. The hydrolysis constant, $K_{(\text{hyd.})}$, of ammonium chloride is

$$0.00000000057$$

Express this value exponentially.

16. The oxidation-reduction constant, $K_{(\text{ox-red.})}$, for the reaction: $5\text{Fe}^{++} + \text{MnO}_4^- + 8\text{H}^+ \rightleftharpoons 5\text{Fe}^{+3} + \text{Mn}^{++} + 4\text{H}_2\text{O}$, is 3.2×10^{63} . How many zeros are required to write out this value of $K_{(\text{ox-red.})}$?

17. Solve:

$$\begin{array}{ll} (a) 4.9 \times 10^2 \times 8.1 \times 10^7; & (d) 2.3 \times 10^8 \times 7.1 \times 10^{-3}; \\ (b) 1.7 \times 10^9 \times 3.5 \times 10^4; & (e) 9.4 \times 10^{-6} \times 6.6 \times 10^6; \\ (c) 9.8 \times 10^5 \times 5.4 \times 10^3; & (f) 4.7 \times 10^{-12} \times 8.3 \times 10^{-11}. \end{array}$$

18. Solve:

$$\begin{array}{ll} (a) \frac{3.2 \times 10^{-14}}{1.9 \times 10^6}; & (b) \frac{8.7 \times 10^{15}}{3.1 \times 10^{12}}; & (c) \frac{1.4 \times 10^2}{9.3 \times 10^{-11}}; & (d) \frac{7.4 \times 10^{-22}}{4.6 \times 10^{-29}}; \\ (e) \frac{2.8 \times 10^3 \times 8.5 \times 10^{-8}}{4.2 \times 10^6}; & (f) \frac{9.6 \times 10^{46} \times 5.1 \times 10^{-24}}{3.2 \times 10^{-15} \times 7.4 \times 10^{33}}. \end{array}$$

19. Convert to correctly expressed exponential numbers:

$$\begin{array}{lll} (a) 391 \times 10^8; & (b) 852.7 \times 10^{-5}; & (c) 16.49 \times 10^{12}; \\ (d) 0.65 \times 10^{-7}; & (e) 0.00034 \times 10^{18}; & (f) 0.0074 \times 10^{-21}. \end{array}$$

20. Solve:

$$\begin{array}{lll} (a) (5.1 \times 10^3)^2; & (b) (1.8 \times 10^{-7})^2; & (c) (8.6 \times 10^4)^3; \\ (d) (0.00072)^3; & (e) (347 \times 10^5)^2; & (f) (0.061 \times 10^{-3})^3. \end{array}$$

21. Solve:

$$\begin{array}{lll} (a) \sqrt{9.4 \times 10^5}; & (b) \sqrt{4.8 \times 10^{-7}}; & (c) \sqrt[3]{1.6 \times 10^{10}}; \\ (d) \sqrt[3]{3.9 \times 10^{-16}}; & (e) (7.2 \times 10^8)^{1/2}; & (f) (2.8 \times 10^{-11})^{1/2}. \end{array}$$

22. Solve:

$$\begin{array}{ll} (a) 7.6 \times 10^6 \times 3.2 \times 10^2 - 1.3 \times 10^9; & \\ (b) \sqrt{3.84 \times 10^{-16} + 5.91 \times 10^{-17} - 2.83 \times 10^{-18}}; & \\ (c) (4.1 \times 10^{-3} \times 9.6 \times 10^{15} - 6.3 \times 10^{21} \times 1.7 \times 10^{-10})^{1/2}. & \end{array}$$

23. Express the following H^+ ion concentrations in terms of pH:

$$(a) 7.3 \times 10^{-5}; \quad (b) 4.6 \times 10^{-3}; \quad (c) 9.5 \times 10^{-12}.$$

24. The $K_{(\text{ion})}$ of NH_4OH is 1.75×10^{-5} . Calculate the $pK_{(\text{ion})}$.

25. The $K_{\text{s.p.}}$ of BaSO_4 is 1.2×10^{-10} . Calculate the $pK_{\text{s.p.}}$.

26. The instability constant of HgBr_4^{2-} is 2.2×10^{-22} . Calculate the $pK_{(\text{instab.})}$.

27. The hydrolysis constant of NH_4Cl is 5.7×10^{-10} . Calculate the $pK_{(\text{hyd.})}$.

28. The oxidation-reduction constant, $K_{(\text{ox-red.})}$, for the reaction $\text{Cl}_2 + 2\text{I}^- \rightleftharpoons 2\text{Cl}^- + \text{I}_2$, is 6.3×10^{27} . Calculate the $pK_{(\text{ox-red.})}$.

29. The $K_{\text{s.p.}}$ of MgNH_4PO_4 is 2.5×10^{-13} . Calculate the Mg^{++} ion concentration in a saturated aqueous solution of MgNH_4PO_4 . *Hint:* $K_{\text{s.p.}} = C_{\text{Mg}^{++}} \times C_{\text{NH}_4^+} \times C_{\text{PO}_4^{3-}}$. In a saturated solution, $C_{\text{Mg}^{++}} = C_{\text{NH}_4^+} = C_{\text{PO}_4^{3-}}$.

30. Solve with the slide rule:

$$(a) \sqrt[6]{4.67 \times 10^{-3} \times 2.47 \times 10^2}$$

Hint: The sixth root is equivalent to the cube root of the square root, or *vice versa*.

$$(b) 10^x = (1.71 \times 4.53)^2.$$

CHAPTER II

SIGNIFICANT FIGURES

In all our researches, indeed, on whatever subject, our object is to arrive at numbers, at quantities, though often in a very imperfect manner and by very uncertain methods.

—AUGUSTE COMTE.

The statement “a chain is no stronger than its weakest link” has become a trite expression. We might paraphrase for our present purpose to “a computed result is no more accurate than the least accurate number entering the computation.” In measuring physical quantities, usually the experimenter knows how finely the scale of the instrument is graduated and what pains have been taken to eliminate likely sources of error. Hence he knows the *degree of exactness* of the measurement. In stating the value of the measured quantity he wants to be honest, and neither exaggerate nor underestimate the exactness. Moreover, the value should be given in such a manner that others can easily obtain that same estimate of the care observed in the experiment. It is therefore necessary to follow a conventional method of stating the result and of carrying out the computations. Beginners in science often needlessly waste both time and energy in making computations. This unnecessary waste can be attributed to the practice of carrying superfluous figures in longhand computations or to the use of five-place logs when only three-place logs are required.

The conventional system that has been universally adopted (but which is often not as universally observed) is incorporated in the following set of rules for significant figures. First let us examine some definitions. **Digits** such as 0, 1, 2, 3, 4 . . . , 9, sometimes called figures, are combined to give numbers, as 12.6 or 0.74. In a number the **significant digits** are those which express as exactly as possible the known magnitudes. To illustrate these definitions: In the number 32.47 (a burette reading) the digit 3 gives the tens, the digit 2 gives the units, the digit 4 gives the

tenths and the digit 7 gives the hundredths. From the graduations on a burette it is easily seen that it is impossible to know the thousandths of a milliliter because the smallest graduation is one-tenth and the value for hundredths is estimated by mentally dividing the smallest division into ten equal parts. In making the above burette reading other observers might have obtained 32.46 or 32.48, but it is highly probable that all observers would report a reading in the range 32.46 to 32.48. The last digit is liable to variation by 1 in either direction and is therefore uncertain, whereas the preceding digits are known with certainty. In any statement of a quantity, the known digits are given and only one uncertain digit is kept. The figures thus written are all significant. The digit 0 requires more consideration since it may or may not be significant. In a burette reading of 30.05 both zeros are significant, and in the weight of a precipitate obtained by the analytical balance, 1.2010 grams, both zeros are significant. If we say that the cost of a battleship is \$50,000,000, none of the zeros is significant. Likewise, in the value 0.0015, which is the solubility of silver chloride in grams per liter of water at 20° C, the zeros merely locate the decimal point. One way of avoiding this difficulty is to state the numbers as exponential numbers, giving all the significant figures in the digit factor. Thus, the cost of a battleship is 5×10^7 dollars, and the solubility of silver chloride is 1.5×10^{-3} gram per liter at 20°. Another way is to choose a larger or smaller unit. Thus 2,500 cm may be expressed as 25 meters, 6,000 pounds may be given as 3 tons, and 0.0251 gram may be written 25.1 mg.

To sum up: the interpretation which should be placed on the statement of a measured quantity is that the last figure given is uncertain by ± 1 . Hence, 1,207 means a value between 1,208 and 1,206; and 6.06×10^{23} denotes a value between 6.07×10^{23} and 6.05×10^{23} .

Rules for Computation:¹

1. *In stating the numerical value of a measured quantity keep only one uncertain figure.*

¹ The rules are, in the main, based upon the discussion of computation rules given by Holman: *Computation Rules and Logarithms*, pp. xii-xv, Macmillan, New York, 1896. Goodwin: *Precision of Measurements*, Second Edition, pp. 21-25, McGraw-Hill, New York, 1920, gives in brief form substantially the rules of Holman.

2. Interpret the uncertainty of the last figure as ± 1 , unless a more precise statement can be given.

3. In dropping superfluous figures, increase the last retained figure by 1 if the first dropped figure (i.e., adjacent to the last retained figure) is 5 or greater. To round off the value 27.0265 to four significant figures the result would be 27.03.

4. In addition and subtraction keep only as many decimals as are given in the number having fewest decimals. To illustrate, find the sum of $13.65 + 0.0082 + 1.632$. As 13.65 has only two decimals, only two need be kept in the other terms. The sum is:

$$\begin{array}{rcll}
 13.65 & \text{is retained as} & 13.65 & \\
 0.0082 & \text{"} & \text{"} & 0.01 \\
 1.632 & \text{"} & \text{"} & 1.63 \\
 \hline
 & & & 15.29
 \end{array}$$

To prove that this is the only possible logical answer, let us take the sum carrying all the decimals and place an x where decimals are unknown.

$$\begin{array}{r}
 13.65xx \\
 0.0082 \\
 1.632x \\
 \hline
 15.29xx
 \end{array}$$

To show this in another way let us take the sums of the maximum values and of the minimum values according to rule 2.

Maximum	Minimum
13.66	13.64
0.0083	0.0081
1.633	1.631
15.3013	15.2791

The sum of the values as stated is:

$$\begin{array}{r}
 13.65 \\
 0.0082 \\
 1.632 \\
 \hline
 15.2902
 \end{array}$$

The sum 15.29 obtained three ways is seen to fall in the range 15.301 to 15.279, and hence the simplest method of addition gives a result that complies with our criteria.

5. *In multiplication and division retain a number of figures that will produce in each factor a percentage uncertainty no greater than that in the factor with fewest significant figures.* For example, the product of the three numbers: $13.65 \times 0.0082 \times 1.632$, should be rewritten $13.7 \times 0.0082 \times 1.63$. The middle factor 0.0082 has only two significant figures, hence the other factors must retain two or, perhaps, three. The uncertainty of this factor is 1 part in 82, roughly 1 in 100 or 1 per cent. Therefore, in order to introduce no greater uncertainty than 1 per cent, three figures must be retained in the other factors. The answer must be stated with such a number of significant figures that the uncertainty in the answer will be no greater than that in the least certain factor entering the computation. Thus, $13.7 \times 0.0082 \times 1.63 = 0.183$. In another example:

$$\begin{array}{r} 1.627 \times 0.0148 \\ \hline 2.3026 \end{array}$$

the factors should be rounded off to

$$\begin{array}{r} 1.63 \times 0.0148 \\ \hline 2.30 \end{array}$$

in accordance with all the rules.

6. *In multiplications and divisions where accuracy no greater than 0.25 per cent is required, the 10-inch slide rule can be used.* Stated roughly, this means that computations involving only three significant figures may usually be performed with a 10-inch slide rule. Note, however, that in this multiplication

$$0.958 \times 9.86$$

the uncertainty is about 1 in 1,000, or 0.1 per cent.

7. *For an accuracy greater than 0.25 per cent, longhand methods or logs must be used in computations.*

8. *In computing with logs, retain in the mantissa of the log only the number of figures that there are significant digits in the digit factors of the numbers entering the computation.* The characteristic denoting, as it does, merely the magnitude cannot be reckoned as significant.

A problem that must be especially noted is calculating pH or related functions which are logarithmic expressions. The number of significant figures in the value for the H^+ ion concentration will dictate the number of decimal places carried out in the pH calculation.

9. *If four or more values are averaged together, an extra figure may be carried in the average.*

10. *In a precision measure, retain only two significant figures.*

The reasons for rules 9 and 10 need not now concern the reader. The discussion of precision measures and averages in Chapter X will clarify them. It is only for convenience in later reference that rules 9 and 10 are included at this point.

More Complex Cases. Examples often arise in which several arithmetical processes are involved and a single computation rule will not suffice. In the solution of

$$k = \left[\frac{0.427(72.6 + 4.38)}{323.7 - 319.3} \right]^2$$

addition, subtraction, multiplication, division, and evolution must be performed. It will be best to carry out the addition and subtraction first, according to rule 4. The value 4.38 will be rounded to 4.4, and the expression becomes

$$k = \left(\frac{0.427 \times 77.0}{4.4} \right)^2$$

Before performing the multiplication and division, all the values will be rounded to two figures, according to rule 5. Then

$$k = \left(\frac{0.43 \times 77}{4.4} \right)^2$$

Slide-rule computation gives

$$k = 57$$

Solving this example has required the successive application of rules 4, 3, 5, 3, and 6.

PROBLEM SET VI

1. Restate these values in more convenient units to eliminate non-significant digits:

- (a) 275,000 grams; (b) 0.271 meter; (c) 0.02 dollar; (d) 0.0064 liter; (e) 17,000 milligrams.

2. Round off these numbers to three figures:

(a) 1.0751; (b) 0.86249; (c) 27.052; (d) 8.971×10^{-16} ; (e) 3.14159.

3. Solve:

$$0.00819 + (4.8 \times 2.12 \times 10^{-4}) - 0.00052008$$

Should slide rule or logs be used?

4. Solve:

$$(a) \frac{48.69 \times 1.86}{95.6 \times 0.050947}; \quad (b) \sqrt[3]{\frac{14.7 + 16.3}{29.8 - 23.5}}$$

5. Given the following atomic weights; calculate the grams of $\text{HC}_2\text{H}_3\text{O}_2$ contained in a liter of 0.231 *M* solution: H = 1.0081; C = 12.010; O = 16.0000.

6. Analysis of a glass sand gave the following values for the percentage of Fe_2O_3 : 0.077, 0.070, 0.079, 0.080, 0.072. Compute the average percentage of Fe_2O_3 according to rule 9.

7. The energy (expressed in ergs) of a quantum of radiation is given by $E = h\nu$; where h is Planck's universal constant with a value of 6.55×10^{-27} erg second, and ν is the frequency of the radiation. The frequency is given by the velocity of light, equal to 3.0×10^{10} cm per sec, divided by the wave length expressed in centimeters. The wave length of a certain blue line in the cadmium spectrum is 4800 Å (1 Å = 10^{-8} cm). Calculate the energy of a quantum of this blue light.

8. The American inch equals 2.53998 centimeters. Express 15.0 cm as inches to the nearest $1/32$ of an inch.

9. Archimedes (287–212 B.C.) proved geometrically that the value of π lay between $3\frac{1}{7}$ and $3\frac{10}{71}$. Ludolph van Ceulen (1539–1610) devoted practically his entire life to the calculation of π and obtained the result, 3.14159265358979323846264338327950288.² If you had measured the diameter of a circle, which was exactly 1 inch, with a meter stick graduated to millimeters, whose value of π would you use to calculate the area of the circle? (a) Perform the calculation. (b) Express the uncertainty of van Ceulen's value of π as an exponential number.

10. The solubility of $\text{Ba}(\text{NO}_3)_2$ is 0.33 mol per liter at 18° C. Calculate the solubility in grams per 100 ml of water, given the atomic weights: Ba = 137.36; N = 14.008; O = 16.0000.

² Ball: *History of Mathematics*, Second Edition, p. 239, Macmillan, London, 1893. It will be noted that Ludolph van Ceulen computed π to thirty-five decimals. In Germany even today π is often called "Ludolph's ratio." The champion of the π -computers is William Shanks, who in 1873 published (*Proc. Roy. Soc. (London)*, 21, 318) a value carried to 707 decimal places. His accuracy cannot be vouched for, since no successor has arisen with sufficient patience to check his computations. However, Shanks does mention that Richter found π to 500 decimals in the year 1853—all of which agree with his published early in the same year.

CHAPTER III

ALGEBRAIC SOLUTION OF NUMERICAL EQUATIONS

Measurement today usually signifies a set of operations performed with measuring rods or other instruments upon a given set of entities to render possible their approximate identification with the set of number-values constituting the scientific object, which in turn can be substituted for the variables in an equation. This is accomplished in proportion as qualitative determinations can be reduced to quantitative ones.

—LEWIS M. HAMMOND.

LINEAR EQUATIONS

The simplest form of numerical equation the chemist is called upon to solve is the linear equation. The simplest linear equation is one in which a single variable occurs only to the first power. A linear equation in general form is

$$mx = b \quad (1)$$

The solution of this equation is

$$x = \frac{b}{m}$$

Or, if

$$2x = 6 \quad (2)$$

Then

$$x = \frac{6}{2} = 3$$

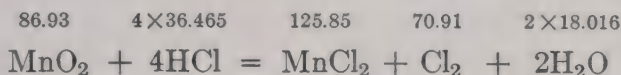
Stoichiometrical Computations. If for a given chemical reaction it is possible to write a balanced chemical equation, the implicitly stated numerical information can be made the basis of stoichiometrical computations. In writing the equation



we are explicitly stating that one mol (gram-molecular-weight) of manganese dioxide reacts with four mols of hydrochloric acid to

produce one mol of manganous chloride, one mol of elementary chlorine, and two mols of water. At the same time we are implicitly stating that 86.93 g of MnO_2 will completely react with 145.86 g of HCl to form 125.85 g of MnCl_2 , 70.91 g of Cl_2 , and 36.03 g of H_2O . Since the sum of the weights of the reactants equals the sum of the weights of the products, another implicit statement is being made to the effect that the law of conservation of mass is obeyed.

It is customary to indicate the weights of the substances involved in the reaction in terms of the molecular weights, thus:



Suppose that it is desired to compute the weight of Cl_2 that would be formed by 8.693 g of MnO_2 . One mol of MnO_2 yields one mol of Cl_2 . It is readily seen that 8.693 g is one-tenth of a mol of MnO_2 . Hence one-tenth of a mol of Cl_2 would be formed. In detail the computation would be

$$\frac{8.693}{86.93} = \frac{x}{70.91} \quad (3)$$

from which

$$x = \frac{8.693}{86.93} \times 70.91 = 7.091$$

In equation (3) the equality between the ratios is an equality between the numbers of mols. The mol or a simple, integral multiple of a mol is the unit of matter participating in a chemical reaction; hence the computation is simplest when reduced to an equality between corresponding basic units.

In fruit commerce, the bushel is the usual unit for measuring goods. A hypothetical trade transaction comparable to the above reaction can be imagined as follows: if a bushel of apples consists of 125 apples and a bushel of peaches consists of 200 peaches, suppose that a dealer exchanges 25 apples for an equal measure of peaches. The result will be

$$\frac{25}{125} = \frac{x}{200} \quad (4)$$

from which

$$x = \frac{25}{125} \times 200 = 40$$

By the exchange, which can be likened to the chemical reaction, 25 apples yield 40 peaches.

Returning to the chemical reaction of MnO_2 with HCl , we see that one mol of MnO_2 yields two mols of H_2O . The unit of water in this reaction is not a single mol but two mols. Hence the weight of water yielded by 8.693 g of MnO_2 is given by

$$\frac{8.693}{86.93} = \frac{x}{2 \times 18.016} \quad (5)$$

from which

$$x = \frac{8.693}{86.93} \times 2 \times 18.016 = 3.603$$

The preceding calculations might have been performed as readily by employing ratio and proportion. In the case just considered the expression

$$8.693 : 86.93 :: x : 2 \times 18.016 \quad (6)$$

might have been written. The same answer would be obtained, but the simplicity of the statement of the reasoning would be lost. In stoichiometrical computations, the process involving the more straightforward line of reasoning from a chemical standpoint is to be preferred. In stating the equality of ratios as compared with stating a proportion, the mathematical processes are equivalent, but the former method is usually preferred by the thoughtful chemist.

Simultaneous Linear Equations. A linear equation may contain more than one unknown as long as the equation is of first degree. The *degree* of an equation is obtained from a term-by-term examination of the equation. In each term the exponents of all variables are summed algebraically. The term having the greatest "exponent sum" dictates the degree. Thus:

$$2x + xy + 5 = 0$$

and

$$x^2 + 3x = 4$$

are both equations of second degree.

If we have two linear equations each involving the same two unknowns, they are simultaneous provided that the two equations express different relationships or are independent.

For example,

$$\begin{aligned} x - 3y &= -1 \\ 4x + 2y &= 10 \end{aligned} \tag{7}$$

are simultaneous; whereas

$$\begin{aligned} 2x - 3y &= 3 \\ 4x - 6y &= 6 \end{aligned} \tag{8}$$

or

$$\begin{aligned} 3x - y &= 5 \\ 4y + z &= 2 \end{aligned} \tag{9}$$

are not simultaneous pairs of equations. Equations (8) are not simultaneous because they are not independent; the second equation of the pair is proportional to the first. Neither are equations (9) simultaneous because they do not involve the same two unknowns.

Two simultaneous equations as (7) can yield a unique value of x and a unique value of y as solutions. In general terms, n independent linear equations each of which involves the same n unknowns are simultaneous and they can yield unique values for the n unknowns.

Elimination. The solution of a set of n simultaneous equations consists in the successive steps of eliminating all but one of the unknowns. The result is the production of n new equations each of which is a linear equation in a single variable. Equations (7)

$$\begin{aligned} x - 3y &= -1 \\ 4x + 2y &= 10 \end{aligned}$$

can be solved by the following steps: (1) Multiply all the terms of the first equation by the coefficient of y in the second equation. (2) Multiply all the terms of the second equation by the negative of the coefficient of y in the first equation. (3) Add the equations. The result is

$$\begin{array}{rcl} 2x - 6y & = & -2 \\ 12x + 6y & = & 30 \\ \hline 14x & = & 28 \end{array} \tag{10}$$

(4) Multiply all the terms of the first equation by the negative of the coefficient of x in the second equation. (5) Multiply all the

terms of the second equation by the coefficient of x in the first equation. (6) Add the two equations. The result is

$$\begin{array}{r} 4x - 12y = -4 \\ -4x - 2y = -10 \\ \hline -14y = -14 \end{array} \quad (11)$$

Equations (10) and (11) constitute the two new equations each of which involves a single unknown. These two simple linear equations

$$\begin{array}{r} 14x = 28 \\ -14y = -14 \end{array}$$

yield the solutions

$$x = \frac{28}{14} = 2$$

$$y = \frac{-14}{-14} = 1$$

The operation of solving a pair of simultaneous linear equations may be interpreted in terms of coordinate geometry. The two linear equations represent two straight lines in the x, y plane. The solution of the equations is a determination of the coordinates of the point at which the two lines intersect.

General Method of Elimination. To display the process of elimination more clearly, we should consider a pair of generalized simultaneous equations, that is, literal equations. Such a pair of equations is

$$\begin{array}{r} a_1x + b_1y = k_1 \\ a_2x + b_2y = k_2 \end{array} \quad (12)$$

The six steps in the process of elimination having been performed successively, the results are

$$\begin{array}{r} a_1b_2x + b_1b_2y = b_2k_1 \\ -a_2b_1x - b_1b_2y = -b_1k_2 \\ \hline (a_1b_2 - a_2b_1)x = b_2k_1 - b_1k_2 \end{array}$$

and

$$\begin{array}{r} -a_1a_2x - a_2b_1y = -a_2k_1 \\ a_1a_2x + a_1b_2y = a_1k_2 \\ \hline (a_1b_2 - a_2b_1)y = a_1k_2 - a_2k_1 \end{array}$$

From which

$$x = \frac{b_2k_1 - b_1k_2}{a_1b_2 - a_2b_1} \quad (13)$$

$$y = \frac{a_1k_2 - a_2k_1}{a_1b_2 - a_2b_1} \quad (14)$$

In these two new equations (13) and (14) we note immediately that the denominators of the two expressions are identical. The terms of this common denominator are written in correct alphabetical order. When the subscripts are in correct numerical order, the sign of the term is positive, but where the subscripts are out of correct numerical order, the sign of the term is negative.

DETERMINANTS

The regularities just noted in the general process of elimination can be made the basis of a new notation by means of which simultaneous equations can be solved quickly and readily. The common denominator of equations (13) and (14) can be written

$$a_1b_2 - a_2b_1 = \begin{vmatrix} a_1 & b_1 \\ a_2 & b_2 \end{vmatrix} \quad (15)$$

provided that we define a set of operational rules such that equation (15) is a true equality.

This new notation is called a determinant. It is always a square array of numbers enclosed by two vertical bars. We denote the vertical lines of numbers as *columns* and the horizontal lines as *rows*. Each number in the array is called an *element*. The number of rows (or columns) gives the *order* of the determinant. Thus the determinant

$$\begin{vmatrix} a_1 & b_1 \\ a_2 & b_2 \end{vmatrix}$$

is of second order. The operational rules by which the numerical value of a second-order determinant is found are:

(1) Form all the possible product terms in which there are no duplications of either letters or subscripts, always arranging the letters in alphabetical order.

(2) If the subscripts of a term are in correct numerical order, the sign of the term is positive; if the subscripts are in inverted order, the sign of the term is negative.

Expanding the above determinant

$$a_1b_2 \text{ and } a_2b_1$$

are possible products, while

$$a_1b_1, a_1a_2, a_2b_2 \text{ and } b_1b_2$$

cannot be terms since there are duplications of either letters or subscripts. Applying rule 2, a_1b_2 is positive while a_2b_1 is negative. Hence

$$\begin{vmatrix} a_1 & b_1 \\ a_2 & b_2 \end{vmatrix} = a_1b_2 - a_2b_1 \quad (16)$$

Compare equation (16) with equation (15). The two are identical. Therefore we have established that the prescribed operational rules make a true equality between a second-order determinant and its expanded form.

It is possible to reduce the expansion of second-order deter-

$$\begin{vmatrix} 1 & -3 \\ 4 & 2 \end{vmatrix} = 1 \cdot 2 - 4(-3)$$

minants to a simple mechanical process. In any determinant of second order, the expansion can be indicated by arrows. The

product denoted by the heavy arrow (sloping downward) is positive, and the product denoted by the light arrow (sloping upward) is negative.

Let us now convert equations (13) and (14) to determinant notation.

In the common denominator, the determinant is made up of the coefficients of x and y in the order in which they occur in the simultaneous equations (12). The two numerator determinants differ from the denominator determinant only in that the column of coefficients of the unknown being solved for in each case is replaced by the

$$x = \frac{\begin{vmatrix} k_1 & b_1 \\ k_2 & b_2 \end{vmatrix}}{\begin{vmatrix} a_1 & b_1 \\ a_2 & b_2 \end{vmatrix}} = \frac{k_1b_2 - k_2b_1}{a_1b_2 - a_2b_1}$$

$$y = \frac{\begin{vmatrix} a_1 & k_1 \\ a_2 & k_2 \end{vmatrix}}{\begin{vmatrix} a_1 & b_1 \\ a_2 & b_2 \end{vmatrix}} = \frac{a_1k_2 - a_2k_1}{a_1b_2 - a_2b_1}$$

column of known numbers from the right-hand side of the equations.

Solving Simple Simultaneous Equations by Determinants. For comparison let us solve equations (7) by means of determinants.

$$x - 3y = -1$$

$$4x + 2y = 10$$

The solution for unique values of x and y is

$$x = \frac{\begin{vmatrix} -1 & -3 \\ 10 & 2 \end{vmatrix}}{\begin{vmatrix} 1 & -3 \\ 4 & 2 \end{vmatrix}} = \frac{(-1)2 - 10(-3)}{1 \cdot 2 - 4(-3)} = \frac{28}{14} = 2$$

$$y = \frac{\begin{vmatrix} 1 & -1 \\ 4 & 10 \end{vmatrix}}{\begin{vmatrix} 1 & -3 \\ 4 & 2 \end{vmatrix}} = \frac{1 \cdot 10 - 4(-1)}{1 \cdot 2 - 4(-3)} = \frac{14}{14} = 1$$

The answers agree perfectly with those obtained previously by the method of elimination.

Determinants of Higher Order. The general process of elimination applied to three simple simultaneous equations results in a general method for solution of the three equations by means of third-order determinants such as

$$\begin{vmatrix} a_1 & b_1 & c_1 \\ a_2 & b_2 & c_2 \\ a_3 & b_3 & c_3 \end{vmatrix}$$

The operational rules for third- (and higher-) order determinants, except for a generalization of rule 2, are the same as for those of second order. Thus:

(1) Form all the possible product terms in which there are no duplications of either letters or subscripts, always arranging the letters in alphabetical order.

(2) If the number of inversions of the subscripts of a term is zero or even, the sign of the term is positive; if the number of inversions is odd, the sign of the term is negative. By the phrase *number of inversions* is meant the number of exchanges of adjacent

subscripts that would be required to transform the subscripts to correct numerical order. For example, the term $a_2b_1c_3$ requires the exchange of 2 and 1. The number of inversions is one, an odd number, hence the term has a negative sign. Further, the term $a_3b_1c_2$ would require the exchange of 3 and 1 to give $a_1b_3c_2$ which requires the exchange of 3 and 2 to give normal order. The number of inversions is two, an even number, hence the sign of the term is positive.

In the expansion of the above determinant of third order the following products are possible terms

$$a_1b_2c_3, a_3b_1c_2, a_2b_3c_1, a_3b_2c_1, a_2b_1c_3, \text{ and } a_1b_3c_2$$

while none of the other 78 possible products¹ can be terms of the expansion according to rule 1. Counting the number of inversions in the six terms of the expansion, the results in tabular form are

TERM	NUMBER OF INVERSIONS	SIGN
$a_1b_2c_3$	0	+
$a_3b_1c_2$	2	+
$a_2b_3c_1$	2	+
$a_3b_2c_1$	3	-
$a_2b_1c_3$	1	-
$a_1b_3c_2$	1	-

Therefore

$$\begin{vmatrix} a_1 & b_1 & c_1 \\ a_2 & b_2 & c_2 \\ a_3 & b_3 & c_3 \end{vmatrix} = a_1b_2c_3 + a_3b_1c_2 + a_2b_3c_1 - a_3b_2c_1 - a_2b_1c_3 - a_1b_3c_2$$

The process of finding the possible terms from a determinant of third order and counting the inversions is a cumbersome method of expansion. The expansion of any third-order determinant such as

$$\begin{vmatrix} 2 & 1 & 4 \\ 4 & 5 & -2 \\ -2 & 3 & 1 \end{vmatrix}$$

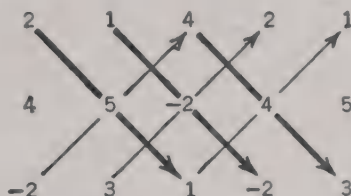
can be reduced to a mechanical process by the following steps:

- (1) Write the determinant omitting the vertical bars.
- (2) Repeat columns 1 and 2 following column 3.

¹ See Problem 7, Chapter IX.

(3) Take the three downward diagonal products (heavy arrows) as positive.

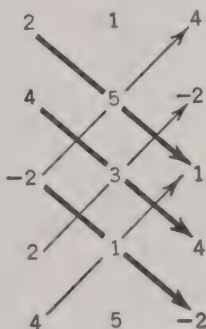
(4) Take the three upward diagonal products (light arrows) as negative. The result is



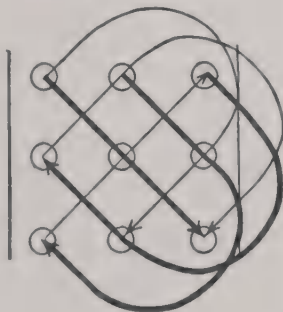
The terms are

$$2 \cdot 5 \cdot 1 + 1(-2)(-2) + 4 \cdot 4 \cdot 3 - (-2)5 \cdot 4 - 3(-2)2 - 1 \cdot 4 \cdot 1 = 110$$

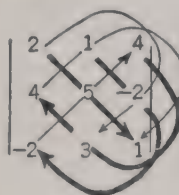
Exactly the same result is obtained by repeating rows instead of columns to give



An alternative method of expanding a third-order determinant depends upon the memorization of a pattern of arrows. This method is quicker and it possesses the advantage of direct operation on the determinant. If the elements are represented by circles, the arrows denote the products in the following diagram.



The heavy arrows denote positive products and the light arrows denote negative products. Using the same numerical determinant as in the preceding paragraphs



$$= 2 \cdot 5 \cdot 1 + 1 \cdot (-2) \cdot (-2) + 4 \cdot 3 \cdot 4 - (-2) \cdot 5 \cdot 4 - 4 \cdot 1 \cdot 1 - 2 \cdot (-2) \cdot 3$$

None of these three mechanical methods of expansion can be applied to determinants of order higher than third. In fact, no practicable mechanical procedures are possible for determinants of fourth, fifth, etc., orders. The expansion of these higher-order determinants is best performed by the method of minors, advantage being taken at the same time of some general properties of determinants to simplify the given determinant. In chemical computations it is hardly ever necessary to deal with more than three simultaneous equations. Therefore we shall not attempt any consideration of these higher-order determinants. Interested students may consult Fine or Sokolnikoff (see bibliography in the Appendix) or any other of the standard textbooks of algebra.

Solution of Three Simple Simultaneous Equations by Determinants. For the sake of simplicity we shall illustrate with an abstract example.

$$3x - y - z = 2$$

$$x - 2y - 3z = -9$$

$$4x + y + 2z = 15$$

The general pattern of the determinants is the same as in the second-order case.

$$x = \frac{\begin{vmatrix} 2 & -1 & -1 \\ -9 & -2 & -3 \\ 15 & 1 & 2 \end{vmatrix}}{\begin{vmatrix} 3 & -1 & -1 \\ 1 & -2 & -3 \\ 4 & 1 & 2 \end{vmatrix}} = \frac{-8 + 45 + 9 - 30 - 18 + 6}{-12 + 12 - 1 - 8 + 2 + 9} = \frac{4}{2} = 2$$

$$y = \frac{\begin{vmatrix} 3 & 2 & -1 \\ 1 & -9 & -3 \\ 4 & 15 & 2 \end{vmatrix}}{2} = \frac{-54 - 24 - 15 - 36 - 4 + 135}{2} = \frac{2}{2} = 1$$

$$z = \frac{\begin{vmatrix} 3 & -1 & 2 \\ 1 & -2 & -9 \\ 4 & 1 & 15 \end{vmatrix}}{2} = \frac{-90 + 36 + 2 + 16 + 15 + 27}{2} = \frac{6}{2} = 3$$

Simultaneous Stoichiometrical Equations. In many instances a sample for quantitative analysis will contain several constituents which are very similar in chemical behavior and hence are not readily separable for individual determination. Recourse is often had to an indirect method of analysis which involves n simultaneous equations when there are n of the similar constituents. The resulting simultaneous equations may be solved readily by means of determinants.

Example. A mixture of AgCl and AgBr weighed 3.1525 g, and the weight of pure Ag obtained from the mixture was 2.0314 g. What were the weights of the AgCl and AgBr in the mixture?

First, two equations must be formulated. Let

$$x = \text{weight of AgCl}$$

$$y = \quad \quad \text{AgBr}$$

$$x + y = 3.1525$$

Also:

$$x \frac{107.88}{143.34} = \text{weight of Ag from } x \text{ g of AgCl}$$

and

$$y \frac{107.88}{187.80} = \text{weight of Ag from } y \text{ g of AgBr}$$

Therefore:

$$x \frac{107.88}{143.34} + y \frac{107.88}{187.80} = 2.0314$$

Which becomes:

$$\begin{aligned} 0.75262x + 0.57444y &= 2.0314 \\ x + y &= 3.1525 \end{aligned}$$

$$x = \frac{\begin{vmatrix} 2.0314 & 0.57444 \\ 3.1525 & 1 \end{vmatrix}}{\begin{vmatrix} 0.75262 & 0.57444 \\ 1 & 1 \end{vmatrix}}; y = \frac{\begin{vmatrix} 0.75262 & 2.0314 \\ 1 & 3.1525 \end{vmatrix}}{\begin{vmatrix} 0.75262 & 0.57444 \\ 1 & 1 \end{vmatrix}}$$

$$x = \frac{2.0314 \times 1 - 3.1525 \times 0.57444}{0.75262 \times 1 - 1 \times 0.57444}$$

$$y = \frac{0.75262 \times 3.1525 - 1 \times 2.0314}{0.75262 \times 1 - 1 \times 0.57444}$$

$$x = \frac{2.0314 - 1.8109}{0.75262 - 0.57444} = \frac{0.2205}{0.17818} = 1.2375$$

$$y = \frac{2.3727 - 2.0314}{0.75262 - 0.57444} = \frac{0.3413}{0.17818} = 1.9155$$

The sum of the calculated weights of the salts (3.1530) does not agree with the observed weight (3.1525). This is another anomaly often observed in evaluating the significance of figures. The result of subtracting two five-figure terms leaves a four-figure term, thereby detracting from the significance of the fourth figure in the result. It is customary, however, to carry five figures throughout the calculation and to regard the error as inherent in the computation of results from indirect methods of chemical analysis.

QUADRATIC EQUATIONS

An equation is termed a quadratic in a single variable when it contains a single unknown raised to a power not higher than second. In other words, the unknown may be present both squared and to the first power. Any quadratic equation in a single variable x may be reduced to the form

$$ax^2 + bx + c = 0$$

and from the equation in just this form, the coefficients a , b , and c are obtained for substitution in the formula we shall derive.

Rearrangement of the preceding equation gives

$$ax^2 + bx = -c$$

now multiply through by $4a$

$$4a^2x^2 + 4abx = -4ac$$

and add b^2 to each side

$$4a^2x^2 + 4abx + b^2 = b^2 - 4ac$$

Inspection reveals that the left side of the equation is a perfect square, or:

$$(2ax + b)^2 = b^2 - 4ac$$

Taking the square root of each side:

$$2ax + b = \pm \sqrt{b^2 - 4ac}$$

Solving for x gives

$$x = \frac{-b \pm \sqrt{b^2 - 4ac}}{2a} \quad (17)$$

This formula shows that there are two possible roots or values of x . To obtain the two roots, separate the equation into the two possible combinations of signs, thus:

$$x_1 = \frac{-b + \sqrt{b^2 - 4ac}}{2a}$$

$$x_2 = \frac{-b - \sqrt{b^2 - 4ac}}{2a}$$

In chemical computations: (1) roots involving the square root of a negative number, *i.e.*, when the term $(b^2 - 4ac)$ has a negative value, could have no significance, and (2) inspection of two real values will usually show immediately which is a possible value and which is an impossible one.

Example. Formic acid ionizes according to the equation:



and the ionization constant is given by

$$K_{(\text{ion})} = \frac{C_{\text{H}^+} \cdot C_{\text{CHO}_2^-}}{C_{\text{HCHO}_2}} = 2.14 \times 10^{-4}$$

What is the H^+ ion concentration in a 0.0120 M solution of formic acid?

Let $x = C_{\text{H}^+} = C_{\text{CHO}_2^-}$. Then

$$2.14 \times 10^{-4} = \frac{x \cdot x}{0.0120 - x}$$

Expressed in the conventional quadratic form

$$x^2 + 2.14 \times 10^{-4}x - 0.0120 \times 2.14 \times 10^{-4} = 0$$

Here

$$a = 1$$

$$b = 2.14 \times 10^{-4}$$

$$c = -0.0120 \times 2.14 \times 10^{-4} = -2.57 \times 10^{-6}$$

$$x = \frac{-b \pm \sqrt{b^2 - 4ac}}{2a}$$

$$x = \frac{-2.14 \times 10^{-4} \pm \sqrt{(2.14 \times 10^{-4})^2 - 4 \times 1(-2.57 \times 10^{-6})}}{2 \times 1}$$

$$x = \frac{-2.14 \times 10^{-4} \pm \sqrt{4.58 \times 10^{-8} + 10.27 \times 10^{-6}}}{2}$$

$$x = \frac{-2.14 \times 10^{-4} \pm \sqrt{10.32 \times 10^{-6}}}{2}$$

$$x = \frac{-2.14 \times 10^{-4} \pm 3.21 \times 10^{-3}}{2}$$

$$x_1 = 1.50 \times 10^{-3}$$

$$x_2 = -1.71 \times 10^{-3}$$

Obviously a negative value for the H^+ ion concentration is meaningless. The correct answer is therefore:

$$C_{H^+} = 1.50 \times 10^{-3} \text{ gram ion per liter}$$

This method of solving the quadratic in a single variable is perfectly general. The student of chemistry need have no misgivings if terms under the square root radical give a negative value. In such cases the answers would be imaginary numbers and the proper conclusion (in a chemical problem) is that he had better check the preceding work, since only real values are answers to the problems that chemists now solve.

EQUATIONS OF HIGHER DEGREE

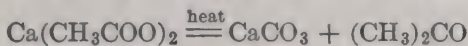
Several general methods of solving the cubic equation have been worked out, and these may be found in almost any advanced textbook of algebra.¹ Algebraic solutions of equations of higher degree than fourth have not been evolved except in certain rare instances. For equations of higher degree than the quadratic, the chemist

¹ Many methods have been proposed for the solution of the general cubic equation, the classical one being that of Cardan. An excellent method is that of Oglesby, *Am. Math. Monthly*, **33**, 321-323(1923).

almost invariably resorts to one of several approximation methods that are capable of yielding answers correct to the desired number of significant figures. Several of the most useful of these approximation methods will be presented in the next chapter.

PROBLEM SET VII²

1. From the reaction expressed by the following equation



calculate the weight of (a) calcium carbonate and (b) acetone yielded by 10.435 g of calcium acetate.

2. Solve for x in the following:

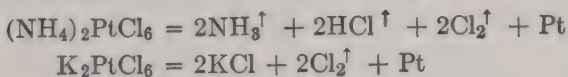
(a) $2.39x + 1.42 - 0.598x = 0.0235$.

(b) $\frac{0.5000}{299.8} = \frac{x}{360.5}$.

(c) $x^2 + 3.68 - 0.529x^2 = 4.05$.

(d) $3.474x - 8.186 = 0$.

3. A solution contained only NH_4Cl and KCl . Upon precipitation with chloroplatinic acid, the precipitate of $(\text{NH}_4)_2\text{PtCl}_6$ and K_2PtCl_6 weighed 0.2141 g. After ignition, the residue of Pt and KCl weighed 0.1433 g. Compute by means of determinants the weights of the two chlorides in the original solution. *Hint:* The reactions involved in the ignition are:



4. Solve by means of determinants:

$$2x + y = 4$$

$$6x + y = 6$$

5. Solve by means of determinants:

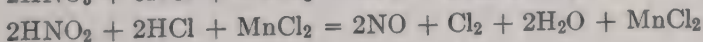
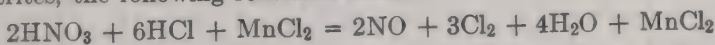
$$x + 2y - 4z = 8$$

$$3x + y + z = 12$$

$$x + y - 5z = 2$$

6. A mixture containing only CaCO_3 and MgCO_3 weighed 0.8739 g. After ignition the resulting mixture of CaO and MgO weighed 0.4510 g. By means of determinants solve for the percentages of the two substances in the original mixture.

7. In the method of Roberts³ for the simultaneous estimation of nitrates and nitrites, the following reactions occur:



² Additional uses for determinants will be found in Chapter XI.

³ *Am. J. Sci.*, [3] 45, 231 (1893).

The liberated chlorine is passed through a solution of potassium iodide, and the iodine set free is determined by titration. The volume of nitric oxide is measured and its weight computed.

The resulting simultaneous equations are:

$$x \frac{3I_2}{2HNO_3} + y \frac{I_2}{2HNO_2} = \text{weight of } I_2$$

$$x \frac{NO}{HNO_3} + y \frac{NO}{HNO_2} = \text{weight of NO}$$

where x is the weight of nitric acid and y is the weight of nitrous acid originally present. In one experiment, a solution containing a mixture of the two acids produced 0.3903 g of iodine and 0.0607 g of nitric oxide. Solve by means of determinants for the weights of nitric and nitrous acids present.

8. In the analysis of gaseous mixtures of C_2H_4 , C_3H_6 , and C_6H_6 , the following equations⁴ must be solved. The values a , b , and c are gas burette readings.

$$x + y + z = a$$

$$2x + 3y + 6z = b$$

$$2x + 2.5y + 2.5z = c$$

By means of determinants, solve for x , y , and z .

9. Solve for x in the following quadratic equations:

$$(a) \quad 1.2x^2 - x - 1.2 = 0.$$

$$(b) \quad -4x^2 + 2x = -30.$$

$$(c) \quad x^2 - 0.08670x + 1.125 \times 10^{-4} = 0.$$

10. For the secondary ionization of sulfuric acid,

$$K_{\text{ion}} = \frac{C_{H^+} \times C_{SO_4^{--}}}{C_{HSO_4^-}} = 3 \times 10^{-2}$$

Calculate the pH of a 0.01 M solution of $NaHSO_4$. Assume the salt to be completely dissociated into Na^+ ions and HSO_4^- ions. *Hint:* To obtain the value of C_{H^+} it is necessary to solve for x in the quadratic,

$$\frac{x \cdot x}{0.01 - x} = 3 \times 10^{-2}$$

⁴ Roscoe and Schorlemmer, *Treatise on Chemistry*, First Edition, Vol. I, p. 704, Macmillan, London, 1878.

CHAPTER IV

APPROXIMATE METHODS OF SOLVING EQUATIONS

Approximation is sometimes defended on the ground that it provides the only knowledge that can be obtained; that the choice is between the approximate knowledge and none at all; and that some knowledge is always better than none.

—N. R. CAMPBELL.

Very often, in the solution of chemical problems, the number of significant figures in the original data does not warrant the labor required for an exact algebraic solution. Then, too, equations are often encountered which are too complex to be solved by the majority of chemists, and there are even equations that the mathematician cannot solve. Resort must then be made to approximate methods.

Method of Neglecting Small Added or Subtracted Quantities.

In problems involving applications of the law of mass action, expressions of this type often arise:

$$\frac{x^2}{m - x} = K$$

in which x is small compared with m . If we assume that x is so small that $m - x$ is approximately equal to m , then the expression may be written

$$\frac{x^2}{m} = K$$

or

$$x = \sqrt{Km}$$

Example 1. For purposes of comparison, let us solve the example on ionization equilibrium used to illustrate the solution of the quadratic equation (page 49). Upon substituting the known values in the general expression we obtained:

$$2.14 \times 10^{-4} = \frac{x^2}{0.0120 - x}$$

54 APPROXIMATE METHODS OF SOLVING EQUATIONS

If we assume x negligible in comparison with 0.0120, the expression becomes

$$2.14 \times 10^{-4} = \frac{x^2}{0.0120}$$

and

$$x = \sqrt{0.0120 \times 2.14 \times 10^{-4}} = \sqrt{2.57 \times 10^{-6}}$$

$$x = 1.60 \times 10^{-3}$$

The answer obtained by the quadratic solution was 1.50×10^{-3} . The approximate method yields an answer which agrees within 6.7 per cent.

$$\frac{1.60 - 1.50}{1.50} \times 100 = 6.7\%$$

Method of Successive Approximations. If a higher degree of accuracy is desired than can be obtained by the preceding method, an extension of the method may be applied. By substituting in the expression the approximate value found by neglecting the small added or subtracted quantity, a closer approximation can be secured. Repetition of the process will eventually produce an answer that agrees exactly with the value from the preceding approximation to the desired number of figures. This answer will be the correct one, and no further manipulation will change the result.

Example 1a. To continue with the example of the last section, let us denote the answer 1.60×10^{-3} as x_1 or the first approximation. Then

$$2.14 \times 10^{-4} = \frac{x_2^2}{0.0120 - x_1} = \frac{x_2^2}{0.0120 - 0.00160}$$

from which

$$x_2^2 = 2.14 \times 10^{-4} \times 0.0104$$

$$x_2 = \sqrt{2.23 \times 10^{-6}}$$

$$x_2 = 1.49 \times 10^{-3}$$

For the third approximation:

$$2.14 \times 10^{-4} = \frac{x_3^2}{0.0120 - 0.00149}$$

$$x_3^2 = 2.14 \times 10^{-4} \times 0.0105$$

$$x_3 = \sqrt{2.25 \times 10^{-6}}$$

$$x_3 = 1.50 \times 10^{-3}$$

As a fourth approximation:

$$\begin{aligned} 2.14 \times 10^{-4} &= \frac{x_4^2}{0.0120 - 0.00150} \\ x_4^2 &= 2.14 \times 10^{-4} \times 0.0105 \\ x_4 &= \sqrt{2.25 \times 10^{-6}} \\ x_4 &= 1.50 \times 10^{-3} \end{aligned}$$

The third and fourth approximations give the same answer, and it is also seen at once that this is the correct answer according to the exact quadratic method of solving.

Often an expression is in such a form that inspection will show whether the first approximation is too large or too small. An according adjustment of the approximate value before solving for the next approximation may save several steps in securing two successive answers that agree. In the foregoing example, in making the first approximation

$$2.14 \times 10^{-4} = \frac{x^2}{0.0120 - x}$$

by neglecting the subtracted x , the answer is obviously too large. It is readily apparent that, if we had diminished the value 1.60×10^{-3} to 1.55×10^{-3} or 1.50×10^{-3} , fewer steps would have been necessary to secure the correct answer.

Method of Graphical Approximation. Graphical treatments of data are, roughly speaking, about as accurate as 10-inch slide-rule computations. When it is desired to solve an equation within this degree of correctness, the equation can be transposed to a general form such as:

$$(1) \quad ax^2 + bx + c = 0$$

or,

$$(2) \quad x^2 + c \log x = 0$$

or, stated more generally,

$$f(x) = 0$$

which means simply that x is the variable in an equation. We can also state

$$y = f(x)$$

meaning that, when a value of x is substituted in the equation, it will become equal to something which we call y . By definition, x

is the independent variable and y is the dependent variable, because the value of y depends upon the random choice of a value of x . Combining all these postulates into one statement, we have:

$$y = f(x) = 0$$

If in the equation a value of x is substituted such that $y = 0$, that value of x is a correct solution of the equation. We can make a graph of the relation $y = f(x)$ by taking successive values of x and solving for corresponding values of y . It is convenient to tabulate the results in a form of this sort:

x	y
x_1	$f(x_1)$
x_2	$f(x_2)$
x_3	$f(x_3)$

Then the graph is prepared by plotting values of y as ordinates and corresponding values of x as abscissas on rectangular (Cartesian) coordinate paper. When the best possible smooth curve is drawn through the points, the roots of the equation are found by locating the values of x at which the curve cuts the x axis. This agrees with the previous statement of the problem because, along the x axis, $y = 0$. The chief advantage of this method is that simple integral values of x may be chosen to make the computations much less tedious. Usually a general knowledge of the physical quantity involved will dictate the proper range of values of x to be chosen.

Example 2. The behavior of gases is closely approximated by van der Waals' equation:

$$\left(p + \frac{a}{V^2}\right)(V - b) = RT$$

For nitrogen the values of the constants a and b are 1.26×10^6 and 47, respectively, when pressure is expressed in atmospheres and volume in cubic centimeters. Let us compute the volume occupied by 1 mol of N_2 at 27°C (300°A) and 80 atmospheres pressure.

First rearrange the equation to the form of a polynomial in V .

$$Y = V^3 - \left(\frac{RT}{p} + b\right)V^2 + \frac{a}{p}V - \frac{ab}{p}$$

Substituting the known values of a , b , p , T , and R (82.07 cc-atm per degree) there is obtained:

$$Y = V^3 - 355V^2 + 15,750V - 740,250$$

For a first approximation let us find what volume the N_2 would occupy if it were a perfect gas. We can calculate from the ideal gas law:

$$pV = RT$$

$$V = \frac{82.07 \times 300}{80} = 308 \text{ cc}$$

As V_1 we shall take 300 and solve for Y .

$$Y_1 = (300)^3 - 355 (300)^2 + 15,750 (300) - 740,250$$

$$Y_1 = -965,250$$

Next try 305 as V_2

$$Y_2 = (305)^3 - 355 (305)^2 + 15,750 (305) - 740,250$$

$$Y_2 = -587,750$$

Taking successive values at intervals of 5, one obtains the following table:

V	Y
300	-965,250
305	-587,750
310	-182,250
315	+251,000
320	+711,750

A graph may now be prepared on rectangular coordinate paper; we shall plot V as abscissas because it is the independent variable, and Y as ordinates because Y depends upon the value of V chosen. By Cartesian convention, positive values of Y will fall above the V axis and negative values below. Fig. 3 shows the plot of the points and the best smooth

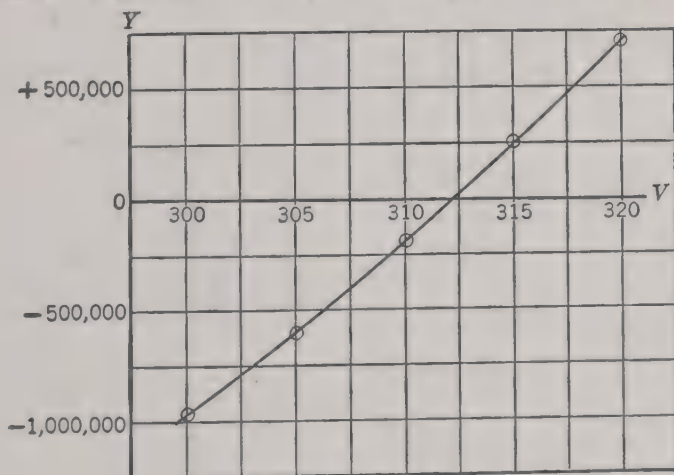


FIG. 3.

curve drawn through the points. The curve cuts the V axis at 312.1. Hence, when $V = 312.1$, $Y = 0$, within the accuracy of plotting.

Newton's Approximation Method. Newton developed a method for solving equations that is applicable to all types of equations, and the accuracy obtainable is limited only by the persistence of the computer. A few exceptional cases cannot be handled by Newton's method, but none of these is likely to be encountered by the chemist. Although the operations are purely mathematical, a graphical interpretation serves best as an introductory explanation. The preceding section has shown how it is possible to plot a graph for a general equation expressed as $y = f(x)$. Any equation of this type can be graphed in the same way.¹ Our

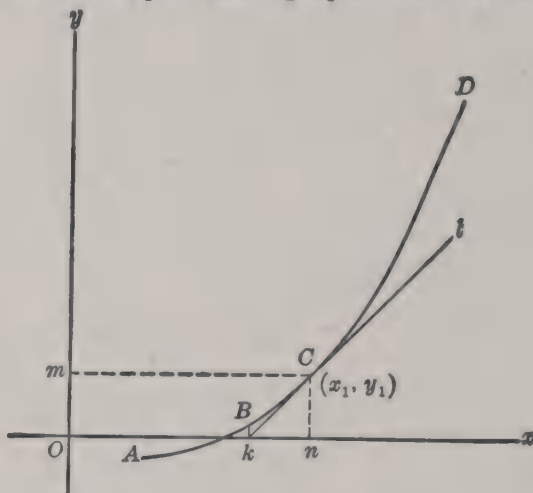


FIG. 4:

object is, as before, to find the value of y most nearly equal to zero. Suppose the equation to be solved gives a graph like $ABCD$ in Fig. 4. Suppose also that the first approximation x_1 (obtained by any one of the several devices already described) gives the value y_1 shown by point C . Next draw a tangent to the curve at C , and extend the tangent to intersect the x axis at k . It is apparent that the length Ok gives a new value x_2 which would correspond to point B on the curve and more nearly approximate the condition of making y equal to zero. Let us then evaluate x_2 in terms of x_1

$$Ok = x_2; On = x_1$$

$$\frac{Cn}{kn} = \tan \angle lkn$$

¹ An excellent treatment of the graphing of the most frequently encountered types of equations is given in Chapters IV, V, and VI, in Daniels: *Mathematical Preparation for Physical Chemistry*, McGraw-Hill, New York, 1928.

From the fundamental definitions of differential calculus the slope of the tangent to a curve at any point is given by dy/dx .

Hence

$$\tan \angle lkn = \left. \frac{dy}{dx} \right|_{at(x_1, y_1)}$$

But

$$Cn = Om = y_1$$

Therefore,

$$kn = \left. \frac{y_1}{\frac{dy}{dx}} \right|_{at(x_1, y_1)}$$

But

$$x_2 = x_1 - kn$$

whence

$$x_2 = x_1 - \left. \frac{y_1}{\frac{dy}{dx}} \right|_{at(x_1, y_1)} \quad (1)$$

The differential coefficient $\frac{dy}{dx}$ is often denoted by $f'(x)$. In keeping with this convention

$$\left. \frac{dy}{dx} \right|_{at(x_1, y_1)} = f'(x_1)$$

Further, our general equation was expressed as

$$y = f(x)$$

from which

$$y_1 = f(x_1)$$

Making these substitutions in equation (1) results in

$$x_2 = x_1 - \frac{f(x_1)}{f'(x_1)} \quad (2)$$

We now have a means of calculating from a first approximation x_1 a better approximation x_2 . A repetition of the process substitut-

ing x_2 for x_1 in the right-hand side of equation (2) gives a still better approximation x_3 .

$$x_3 = x_2 - \frac{f(x_2)}{f'(x_2)}$$

This repetition can be represented graphically by drawing a tangent to the curve at B . The point of intersection of the new tangent with the x axis gives the value of x_3 , which will come even nearer making y equal to zero.

The accuracy that can be obtained by continued repetition is unlimited. In practice, approximations are continued until successive answers agree to the number of significant figures dictated by the original data of the problem.

Newton's method provides a purely analytical procedure for approximating, as closely as is desired, the solution of all types of equations. Not every equation, however, can be treated successfully by this method. If a graph of the function shows a maximum, a minimum, or a point of inflexion near its intersection with the x axis, succeeding approximations will be worse rather than better. It may be reiterated that this type of function is rarely encountered in the mathematics of chemistry.

Example 3. Solve the equation

$$x = 8 \log x$$

to five significant figures.

First, transpose the expression to

$$x - 8 \log x = 0$$

The condition sought is

$$f(x) = x - 8 \log x = 0$$

just as in the example already discussed under "graphical approximation."

As a first approximation, we shall substitute a few simple, integral values of x . Taking $x = 1$,

$$f(x) = 1 - 0 = 1$$

Taking $x = 2$

$$f(x) = 2 - 8 \times 0.30 = -0.4$$

The value of x sought lies between 1 and 2. Compiling a table,

x	$f(x)$
1	1
1.5	0.09
1.6	-0.03
2	-0.4

Therefore we shall choose, as our first approximation, $x_1 = 1.58$.

To apply the correction formula (2), we must obtain $f'(x)$. Since our equation involves a Briggsian log, we must convert to a Napierian log, thus

$$f(x) = x - 8 \times 0.43429 \ln x$$

from which

$$f'(x) = \frac{d}{dx} f(x) = 1 - \frac{8 \times 0.43429}{x}$$

Substituting in equation (2)

$$x_2 = 1.58 - \frac{1.58 - 8 \log 1.58}{1 - \frac{8 \times 0.434}{1.58}}$$

$$x_2 = 1.58 - \frac{1.58 - 1.5893}{1 - 2.2} = 1.58 - 0.0077$$

$$x_2 = 1.5723$$

Repeating the process

$$x_3 = 1.5723 - \frac{1.5723 - 8 \log 1.5723}{1 - \frac{8 \times 0.43429}{1.5723}}$$

$$x_3 = 1.5723 - \frac{1.5723 - 1.5723}{1 - 2.2} = 1.5723 - 0.0000$$

$$x_3 = 1.5723$$

Here it is seen that the first application of Newton's formula using a three-figure approximation gives the correct results to five figures as shown by the exact agreement of x_2 and x_3 . The computer, however, will not always be so fortunate.

Example 4. Young's² empirical equation

$$T' = T + \frac{144.86}{T^{0.0148} \sqrt{T}}$$

expresses the relation between the boiling point (T) of a given homologue and the boiling point (T') of the next higher homologue of the series, where temperatures are expressed in degrees absolute. When a value of T is given from which to compute T' , the problem is straightforward arithmetic. However, given the information that the boiling point of heptane, C_7H_{16} , is 371.30°A , the calculation of the boiling point of hexane,

² S. Young, *Phil. Mag.*, **9**, 1-19 (1905). The equation holds better for saturated hydrocarbons than for any other series.

C_6H_{14} , presents a problem which is not solvable by an algebraic formula. The equation

$$371.30 = T + \frac{144.86}{T^{0.0148} \sqrt{T}}$$

can be solved³ by first transposing to give

$$371.30 - T = \frac{144.86}{T^{0.0148} \sqrt{T}}$$

and taking logs of both sides

$$\log(371.30 - T) = \log 144.86 - 0.0148 \sqrt{T} \log T$$

We next set the condition

$$f(T) = \log(371.30 - T) - \log 144.86 + 0.0148 \sqrt{T} \log T = 0$$

Differentiating, one obtains

$$f'(T) = 0.43429 \left[\frac{-1}{371.30 - T} + 0.0148 \left(\sqrt{T} \frac{1}{T} + \ln T \frac{1}{2\sqrt{T}} \right) \right]$$

which becomes

$$f'(T) = -\frac{0.43429}{371.30 - T} + \frac{0.0148}{\sqrt{T}} \left(0.43429 + \frac{\log T}{2} \right)$$

According to the generalizations of Kopp⁴ and other early physical chemists, we should expect in this case a difference in boiling points of about 30°. Let us take 342° A as a first approximation. Then

$$T_2 = 342 - \frac{\log 29.3 - \log 144.9 + 0.0148 \sqrt{342} \log 342}{-\frac{0.434}{29.3} + \frac{0.0148}{\sqrt{342}} \left(0.434 + \frac{\log 342}{2} \right)}$$

$$T_2 = 342 - \frac{1.4667 - 2.1611 + 0.6936}{-0.1483 + 0.00136} = 341.94$$

Taking 341.94 as a second approximation, the next step gives

$$T_3 = 341.94 - \frac{\log 29.36 - \log 144.86 + 0.0148 \sqrt{341.94} \log 341.94}{-\frac{0.4343}{29.36} + \frac{0.0148}{\sqrt{341.94}} \left(0.4343 + \frac{\log 341.94}{2} \right)}$$

$$T_3 = 341.94 - \frac{1.46776 - 2.16095 + 0.69349}{-0.01479 + 0.00136} = 341.96$$

Further applications of Newton's formula would be expected to produce a final result that differs by not more than ± 1 in the second place of deci-

³ The condition:

$$\phi(T) = T - 371.30 + 144.86 T^{-0.0148} \sqrt{T} = 0$$

could be set and the problem worked using this form. However, the differentiation involved in obtaining $\phi'(T)$ would be much more difficult than with the log equation.

⁴ Smiles: *Chemical Constitution and Physical Properties*, Chapter VII, Longmans, Green and Co., London, 1910.

mals. In any event, a comparison of T_3 with the observed result 341.95°A shows that excellent agreement has been obtained even with this small number of approximations.

PROBLEM SET VIII

1. The calculation of the hydrogen-ion concentration in a 1.00 M solution of HNO_2 which is also 0.050 M with respect to NH_4NO_2 requires the solution of the equation:

$$4.0 \times 10^{-4} = \frac{x(0.050 + x)}{1.00 - x}$$

(assuming complete dissociation of the NH_4NO_2). Solve for x by successive approximations.

2. To a solution containing 0.0025 gram ion per liter of Ti^+ ion, CNS^- ions are added to the extent of 0.48 gram ion per liter. The calculation of the amount of TiCl that must precipitate to establish equilibrium according to the solubility-product principle requires the solution of the equation:

$$(0.0025 - x)(0.48 - x) = 1.4 \times 10^{-4}$$

Solve for x by successive approximations.

3. In a study of the heterogeneous equilibrium involved in the distribution of benzoic acid between water and benzene this equation⁵ occurs,

$$K = \frac{[C_1(1 - \alpha)]^2}{C_2k^2 - C_1k(1 - \alpha)}$$

Given:

$$K = 0.0286; C_1 = 0.0823; C_2 = 0.4726; k = 0.700$$

Calculate α by successive approximations.

4. The equation⁶

$$\log P = 23.8296 - \frac{3,480.3}{T} - 5.081 \log T$$

gives the relation between vapor pressure, P (in millimeters of Hg), and temperature, T (in degrees absolute), for *o*-toluidine. The true boiling point is that temperature at which the vapor pressure is 760 mm. Determine the boiling point of *o*-toluidine by graphical approximation.

5. Solve the following equation for Q by Newton's method:

$$4.52 Q = 1.56 e^Q$$

6. The symbols in the equation⁶ for *m*-toluidine

$$\log P = 18.5043 - \frac{3,200.9}{T} - 3.323 \log T$$

have the same significance as in the equation of Problem 4 above. By Newton's method, solve for the boiling point of *m*-toluidine.

⁵ Taylor and Taylor: *Elementary Physical Chemistry*, Second Edition pp. 378-9. Van Nostrand, New York, 1937.

⁶ Berliner and May, *J. Am. Chem. Soc.*, **49**, 1007-11 (1927).

CHAPTER V

INTERPOLATION AND EXTRAPOLATION

As a science grows more exact it becomes possible to employ more extensively the accurate and concise methods and notation of mathematics.

—G. N. LEWIS and M. RANDALL.

LINEAR INTERPOLATION

It very frequently happens that tables of data or tables of numerical figures are given at fairly large intervals in the values of the independent variable. When it is necessary to know an intermediate value no trouble is involved in finding it by the simple method of similar triangles (sometimes called the rule of proportional parts) as long as the value is a linear function of the

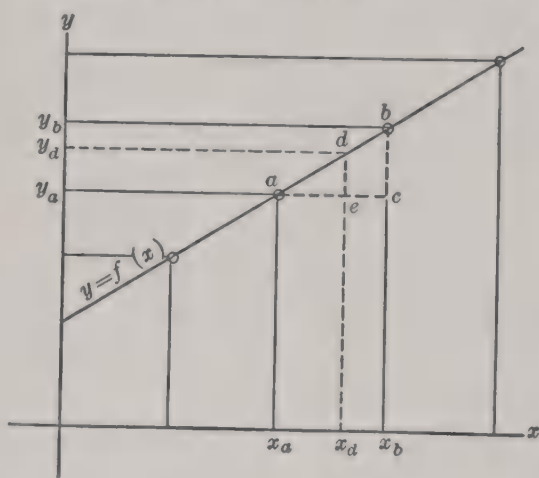


FIG. 5.

independent variable. When the function is not linear, the similar triangle method can yield only a rough approximation. The determination (either arithmetically or graphically) of the required value of a varying quantity, of which certain particular values within the range are already known, is called interpolation.

Graphically the simple linear case is represented by Fig. 5. The triangles abc and ade are similar, hence

$$ac : ae = cb : ed \quad (1)$$

But:

$$ac = x_b - x_a$$

$$ae = x_d - x_a$$

$$cb = y_b - y_a$$

$$ed = y_d - y_a$$

It follows that:

$$\frac{x_b - x_a}{x_d - x_a} = \frac{y_b - y_a}{y_d - y_a} \quad (2)$$

from which

$$y_d = y_a + \frac{(x_d - x_a)}{(x_b - x_a)} (y_b - y_a) \quad (3)$$

Equation (3) is a formula by means of which the value of y corresponding to a value of x between two other values x_a and x_b , for which y_a and y_b are known, can be computed.

To perform an inverse interpolation (i.e., find the value of x when an intermediate value of y is known), rearrange equation (2):

$$x_d = x_a + \frac{(y_d - y_a)}{(y_b - y_a)} (x_b - x_a) \quad (4)$$

Example. A table of the circumferences of circles of given radii is as follows:

RADIUS (r)	CIRCUMFERENCE (C)
1	6.28
2	12.57
3	18.85
4	25.13
5	31.41

Find the circumference of a circle of radius 3.3.

$$C_{3.3} = C_3 + \frac{(r_{3.3} - r_3)}{(r_4 - r_3)} (C_4 - C_3)$$

$$C_{3.3} = 18.85 + \frac{(3.3 - 3.0)}{(4.0 - 3.0)} (25.13 - 18.85)$$

$$C_{3.3} = 18.85 + 0.3(6.28)$$

$$C_{3.3} = 18.85 + 1.88 = 20.73$$

To illustrate inverse interpolation, let us find the radius of a circle the circumference of which is 15.71.

$$r_{15.71} = r_2 + \frac{(C_{15.71} - C_{12.57})}{(C_{18.85} - C_{12.57})} (r_3 - r_2)$$

$$r_{15.71} = 2.0 + \frac{15.71 - 12.57}{18.85 - 12.57} (3.0 - 2.0)$$

$$r_{15.71} = 2.0 + \frac{3.14}{6.28} (1.0)$$

$$r_{15.71} = 2.5$$

We know that the circumference of a circle is a linear function of the radius.

$$C = 2\pi r$$

shows that the equation is of the general type for a straight-line function:

$$y = b + mx$$

Here the constant b equals zero and the slope m equals 2π .

NON-LINEAR INTERPOLATION

When the function is not a linear one it cannot be represented by a simple equation of the straight-line type. Fig. 6 shows the

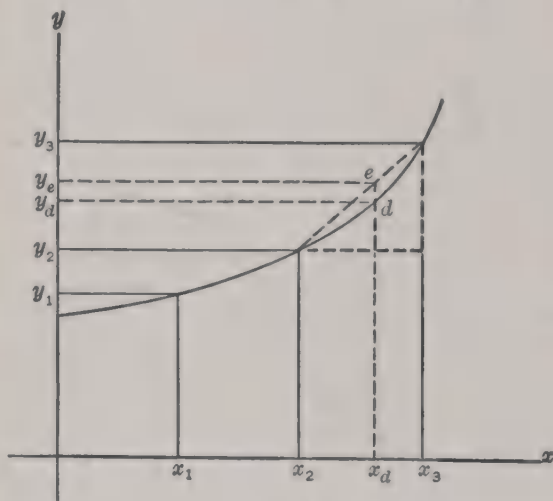


FIG. 6.

error that would be introduced if the rule of proportional parts were used for interpolation. Instead of the correct value y_d corresponding to x_d , the erroneous value y_e would be obtained by the method of similar triangles. Under such conditions the values of y might be plotted against corresponding values of x and intermediate values read

directly from the graph. The accuracy attainable with graphical methods is limited, and oftentimes the chemist has not at hand the materials for such a graphical construction. It will be well, therefore, to consider an analytical method of interpolation for functions of all degrees.

Polynomials. Our general interpolation method is due to Newton, and it is based upon the elementary principles of the

*calculus of finite differences.*¹ First we must consider a demonstrated fact of mathematics, namely, *almost any function may be represented to a fair degree of accuracy by a polynomial equation of the type:*

$$U_x = f(x) = a + bx + cx^2 + dx^3 + \cdots + mx^n \quad (5)$$

For example, the coefficient of linear expansion of a substance is usually represented by

$$l = a + bt$$

Also, specific heats are represented by such equations as

$$C_p = \alpha + \beta T + \gamma T^2$$

When data have been tabulated, it usually means either that the exact theoretical form of the function is unknown or that, if known, it is too complex for ready computation. Lacking better information, we have to assume that the data under consideration can be represented within the required accuracy by a polynomial. The highest power to which the independent variable appears gives the degree of the polynomial. The equation

$$U_x = a - cx^2 + dx^3$$

is a third-degree polynomial.

Differences. If the data are given with a constant increment between successive values of the independent variable, we may set up a table and take differences between successive pairs of values of the function.

¹ The word "calculus" denotes a systematized branch of mathematics and is not restricted to the mathematics of infinitesimals or of integrals. Its source is the Latin *calculi*, which means "stones" or "pebbles." The name arose from the use of pebbles on the counting boards (*abaci*) on which the early Roman arithmetic was performed. There are many varieties of calculus, such as differential calculus, integral calculus, calculus of variations, calculus of finite differences, calculus of operations, and vector calculus.

It was Sir Isaac Newton (1642-1727), regarded as foremost among mathematicians and scientists of all time, who evolved the methods and basic equations of the calculus of finite differences. It remained, however, for a pupil of Newton, Brook Taylor (1685-1731), to systematize and to establish the rigorous proofs of this calculus. The first comprehensive treatise on the subject was published in 1717, *Traite du calcul des différences finies*, by François Nicole (1683-1758). For more recent works consult the Bibliography in the Appendix.

x	U_x	ΔU_x	$\Delta^2 U_x$	$\Delta^3 U_x$
0	0			
		$1 - 0 = 1$		
1	1		$3 - 1 = 2$	
		$4 - 1 = 3$		$2 - 2 = 0$
2	4		$5 - 3 = 2$	
		$9 - 4 = 5$		$2 - 2 = 0$
3	9		$7 - 5 = 2$	
		$16 - 9 = 7$		
4	16			

Δ means the difference between a pair of values, Δ^2 is the same as $\Delta(\Delta)$ or the difference between successive differences, etc. Here it is seen that the column of second differences (Δ^2) is constant and the third differences are therefore zero. The preceding table consists of values of $U_x = x^2$, that is to say, it is a second-degree equation.

Consider next a table of values for $U_x = 2 + x + x^3$

x	U_x	ΔU_x	$\Delta^2 U_x$	$\Delta^3 U_x$	$\Delta^4 U_x$
0	2				
		$4 - 2 = 2$			
1	4		$8 - 2 = 6$		
		$12 - 4 = 8$		$12 - 6 = 6$	
2	12		$20 - 8 = 12$		$6 - 6 = 0$
		$32 - 12 = 20$		$18 - 12 = 6$	
3	32		$38 - 20 = 18$		$6 - 6 = 0$
		$70 - 32 = 38$		$24 - 18 = 6$	
4	70		$62 - 38 = 24$		$6 - 6 = 0$
		$132 - 70 = 62$		$30 - 24 = 6$	
5	132		$92 - 62 = 30$		
		$224 - 132 = 92$			
6	224				

This time the third difference is constant, and the polynomial is of the third degree.

Now let us take the table of circumferences and radii from the example on page 65:

r	C	ΔC	$\Delta^2 C$
1	6.28		
		6.29	
2	12.57		0.01
		6.28	
3	18.85		0.00
		6.28	
4	25.13		

The column of first differences is constant (Rule 2 for significant figures), and this is an example of a first-degree polynomial. Upon the basis of the three cases considered, the conclusion might be drawn that the degree of the polynomial dictates which difference will be constant. Actually, this conclusion is entirely true and is readily susceptible of proof with a polynomial of any degree. In all the preceding tables, the independent variable is given at intervals of 1 unit. The interval may be any value (h) so long as it is constant. A generalized difference table will appear thus:

x	U_x	ΔU_x	$\Delta^2 U_x$	$\Delta^3 U_x$
a	U_a			
		ΔU_a		
$a + h$	U_{a+h}		$\Delta^2 U_a$	
		ΔU_{a+h}		$\Delta^3 U_a$
$a + 2h$	U_{a+2h}		$\Delta^2 U_{a+h}$	
		ΔU_{a+2h}		
$a + 3h$	U_{a+3h}			

From the table it is seen that

$$\Delta U_a = U_{a+h} - U_a$$

and

$$\Delta^2 U_a = \Delta U_{a+h} - \Delta U_a$$

or in general

$$\Delta^r U_a = \Delta^{r-1} U_{a+h} - \Delta^{r-1} U_a \quad (6)$$

Operators. We must introduce a new symbol which represents the operation of stepping by the interval h in the column of values of U_x

$$EU_a = U_{a+h}$$

But since

$$\Delta U_a = U_{a+h} - U_a$$

$$U_{a+h} = U_a + \Delta U_a$$

or

$$EU_a = U_a + \Delta U_a$$

E and Δ are symbols denoting an operation, whereas U is a symbol of function. Symbols of operation, called **operators**, may be classified into two divisions, namely, separable and non-separable. For an operator to be separable it must obey three laws of algebra. These laws are:

$$(1) \quad a(b + c) = ab + ac$$

Distributive law

$$(2) \quad ab = ba$$

Commutative law

$$(3) \quad a^m \cdot a^n = a^{m+n}$$

Exponent summation law

Both Δ and E obey all three laws and hence are separable. For example,

$$\begin{aligned}\Delta(U_a + V_a) &= (U_{a+h} + V_{a+h}) - (U_a + V_a) \\ &= (U_{a+h} - U_a) + (V_{a+h} - V_a)\end{aligned}$$

$$\Delta(U_a + V_a) = \Delta U_a + \Delta V_a$$

Δ therefore obeys the distributive law.

Also,

$$\Delta(kU_a) = kU_{a+h} - kU_a = k(U_{a+h} - U_a)$$

$$\Delta(kU_a) = k(\Delta U_a)$$

displaying obedience to the commutative law. Further,

$$\Delta(\Delta U_a) = \Delta(U_{a+h} - U_a)$$

which, as has already been proved, becomes

$$\Delta(\Delta U_a) = \Delta U_{a+h} - \Delta U_a$$

By definition,

$$\Delta^2 U_a = \Delta U_{a+h} - \Delta U_a$$

Therefore,

$$\Delta(\Delta U_a) = \Delta^2 U_a$$

and the law of summation of exponents is obeyed. These proofs may be extended to other examples, but this extension, as well as proof that E obeys the same laws, will be left to the student.

Examples of non-separable operators are

$$\log(x + y) \neq \log x + \log y$$

because

$$\log x + \log y = \log(xy)$$

Also

$$\sin(x + y) \neq \sin x + \sin y$$

because,

$$\sin(x + y) = \sin x \cdot \cos y + \cos x \cdot \sin y$$

Difference Formula of Newton. Since the operators Δ and E are separable, they can be factored just as any algebraic symbol can. Returning to the expression

$$\begin{aligned}
 EU_a &= U_a + \Delta U_a \\
 (E)U_a &= (1 + \Delta)U_a \\
 E &\equiv 1 + \Delta
 \end{aligned}
 \tag{7}$$

For a general formulation

$$U_{a+mh} = E^m U_a$$

But

$$E^m = (1 + \Delta)^m \tag{8}$$

By the binomial theorem

$$(1 + \Delta)^m = 1 + m\Delta + \frac{m(m-1)}{2!} \Delta^2 + \frac{m(m-1)(m-2)}{3!} \Delta^3 + \dots$$

The symbol of function may now be recombined with the expanded operator

$$U_a \left(1 + m\Delta + \frac{m(m-1)}{2!} \Delta^2 + \frac{m(m-1)(m-2)}{3!} \Delta^3 + \dots \right)$$

Or

$$\begin{aligned}
 U_{a+mh} &= U_a + m\Delta U_a + \frac{m(m-1)}{2!} \Delta^2 U_a \\
 &\quad + \frac{m(m-1)(m-2)}{3!} \Delta^3 U_a + \dots
 \end{aligned}
 \tag{9}$$

This expression is Newton's interpolation formula in its most general form. In order to apply the formula to any specific problem, the value of m must first be evaluated. If we want to solve for the value of U corresponding to a given value of x , we have

$$U_x = U_{a+mh}$$

from which

$$x = a + mh$$

and

$$m = \frac{x - a}{h} \tag{10}$$

It will be well to recapitulate the exact meanings of the symbols: x is the value of the independent variable for which the value of U is sought; a is the value of the independent variable in the first term of the table; and h is the value of the *constant increment* in the independent variable.

Example 1.

x	U_x	Δ	Δ^2	Δ^3
2	1.2			
		9.0		
4	10.2		8.0	
		17.0		0.0
6	27.2		8.0	
		25.0		0.0
8	52.2		8.0	
		33.0		
10	85.2			

Solve for U when $x = 5$.

$$m = \frac{x - a}{h} = \frac{5 - 2}{2} = 1.5$$

$$U_5 = U_2 + m\Delta U_2 + \frac{m(m-1)}{2 \cdot 1} \Delta^2 U_2$$

$$U_5 = 1.2 + 1.5 \cdot 9.0 + \frac{1.5 \cdot 0.5}{2 \cdot 1} \cdot 8.0$$

$$U_5 = 1.2 + 13.5 + 3.0 = 17.7$$

Proof:

The above function is

$$U_x = 0.2 - 1.5x + x^2$$

$$U_5 = 0.2 - 7.5 + 25.0 = 17.7$$

Example 2. As a less abstract case, the following table taken from the data of Bingham and Jackson, *Bulletin, Bureau of Standards*, **14**, 75 (1917), gives the viscosities of water (expressed in centipoises) over a 10-degree temperature range:

$t^\circ \text{C}$	μ	Δ	Δ^2	Δ^3	Δ^4
20	1.0050				
		-.0471			
22	0.9579		+.0034		
		-.0437		-.0002	
24	0.9142		+.0032		-.0002
		-.0405		-.0004	
26	0.8737		+.0028		.0000
		-.0377		-.0004	
28	0.8360		+.0024		
		-.0353			
30	0.8007				

Solve for μ when $t = 25^\circ \text{C}$

$$m = \frac{x - a}{h} = \frac{25 - 20}{2} = 2.5$$

$$\mu_{25} = \mu_{20} + 2.5\Delta\mu_{20} + \frac{2.5(2.5-1)}{2 \cdot 1} \Delta^2\mu_{20} + \frac{2.5(2.5-1)(2.5-2)}{3 \cdot 2 \cdot 1} \Delta^3\mu_{20}$$

$$\mu_{25} = 1.0050 + 2.5(-0.0471) + 1.88(0.0034) + 0.313(-0.0002)$$

$$\mu_{25} = 0.8936$$

This example reveals a more realistic case than Example 1. Here it is seen that exact constancy of one of the columns of differences is not attained. It is to be noted that the Δ^3 column is most nearly constant of any of the columns of differences. In the Δ^4 column the sign begins to oscillate from $-$ to $+$, and from there on the differences would diverge rapidly because of the changes in sign.

In all cases it is necessary to use judgment in choosing the column which is most nearly constant. A possible criterion is to calculate the value of one extra term in the interpolation formula and see whether it affects the significant figures of the answer.

Variable Increments. Quite often data are given in a form in which the interval between successive values of the independent variable is not constant. Two methods are available between which there is little choice. These will be considered in the order: (1) Newton's divided difference formula, and (2) Lagrange's interpolation formula.

Newton's Divided Difference Formula. This method is an extension of the foregoing method for constant intervals. Rigorous proof will not be attempted, but the assurance can be given that the principles are the same except that they are a little more generalized. A tabulation of the divided differences (Δ') for a function of x will be as follows:

x	U_x	Δ'	Δ'^2	Δ'^3
a	U_a	$\frac{U_b - U_a}{b - a}$		
b	U_b	$\frac{U_c - U_b}{c - b}$	$\frac{\Delta'U_b - \Delta'U_a}{c - a}$	$\frac{\Delta'^2U_b - \Delta'^2U_a}{d - a}$
c	U_c	$\frac{U_d - U_c}{d - c}$	$\frac{\Delta'U_c - \Delta'U_b}{d - b}$	$\frac{\Delta'^2U_c - \Delta'^2U_b}{e - b}$
d	U_d	$\frac{U_e - U_d}{e - d}$	$\frac{\Delta'U_d - \Delta'U_c}{e - c}$	
e	U_e			

The interpolation formula in general form is:

$$U_x = U_a + (x - a)\Delta'U_a + (x - a)(x - b)\Delta'^2U_a \\ + (x - a)(x - b)(x - c)\Delta'^3U_a + \dots \quad (11)$$

Just as before, the degree of the polynomial determines which column of divided differences will be constant, and the formula is used carrying the number of terms dictated by which difference is constant, or most nearly so.

Example 3.

x	U_x	Δ'	Δ'^2	Δ'^3
1	0			
		$\frac{4 - 0}{2 - 1} = 4$		
2	4		$\frac{32 - 4}{5 - 1} = 7$	
		$\frac{100 - 4}{5 - 2} = 32$		$\frac{13 - 7}{7 - 1} = 1$
5	100		$\frac{97 - 32}{7 - 2} = 13$	
		$\frac{294 - 100}{7 - 5} = 97$		$\frac{19 - 13}{8 - 2} = 1$
7	294		$\frac{154 - 97}{8 - 5} = 19$	
		$\frac{448 - 294}{8 - 7} = 154$		
8	448			

Find U_4 .

$$U_4 = U_1 + (4 - 1)\Delta'U_1 + (4 - 1)(4 - 2)\Delta'^2U_1 + \\ (4 - 1)(4 - 2)(4 - 5)\Delta'^3U_1$$

$$U_4 = 0 + 3 \cdot 4 + 3 \cdot 2 \cdot 7 + 3 \cdot 2 \cdot (-1)1$$

$$U_4 = 12 + 42 - 6$$

$$U_4 = 48$$

Proof:

The function is

$$U_x = -x^2 + x^8$$

$$U_4 = -16 + 64 = 48$$

Example 4. The following values of the molar heat capacity of ethylene are taken from the data of Egan and Kemp, *J. Am. Chem. Soc.*, **59**, 1264-8 (1937). From these values we shall compute C_p at 55.00°A by means of Newton's method of divided differences.

$T^\circ \text{A}$	C_p	Δ'	Δ'^2	Δ'^3
42.08	6.045			
		$\frac{0.887}{5.21} = 0.170$		
47.29	6.932		$\frac{-0.002}{10.44} = -0.0002$	
		$\frac{0.877}{5.23} = 0.168$		$\frac{-0.0011}{15.52} = -0.00007$
52.52	7.809		$\frac{-0.013}{10.31} = -0.0013$	
		$\frac{0.789}{5.08} = 0.155$		$\frac{-0.0011}{15.51} = -0.00007$
57.60	8.598		$\frac{-0.024}{10.28} = -0.0024$	
		$\frac{0.681}{5.20} = 0.131$		
62.80	9.279			

$$\begin{aligned}
 C_{p55} &= 6.045 + (55.00 - 42.08) 0.170 + (55.00 - 42.08)(55.00 - 47.29)(-0.0002) \\
 &\quad + (55.00 - 42.08)(55.00 - 47.29)(55.00 - 52.52)(-0.00007) \\
 &= 6.045 + 2.193 - 0.020 - 0.017
 \end{aligned}$$

$$C_{p55} = 8.201$$

Lagrange's Interpolation Formula. Just as in the Newton method, the assumption must first be made that a function may be represented by a polynomial expression. Specifically, Lagrange assumed that the function for which n values are given with any interval in the independent variable may be represented by an

expression of n terms, each term being of the $(n - 1)$ th degree in x .

$$\begin{aligned} U_x = & A(x - b)(x - c)(x - d) \cdots (x - n) \\ & + B(x - a)(x - c)(x - d) \cdots (x - n) \\ & + C(x - a)(x - b)(x - d) \cdots (x - n) \\ & + D(x - a)(x - b)(x - c) \cdots (x - n) + \cdots \end{aligned}$$

Let us now evaluate $A, B, C, D \cdots$; when $x = a$

$$U_a = A(a - b)(a - c)(a - d) \cdots (a - n)$$

All the other terms become zero.

Therefore:

$$A = \frac{U_a}{(a - b)(a - c)(a - d) \cdots (a - n)}$$

Likewise

$$B = \frac{U_b}{(b - a)(b - c)(b - d) \cdots (b - n)}$$

And the final form of the expression when the values of $A, B, C, D \cdots$ are substituted is

$$\begin{aligned} U_x = & U_a \frac{(x - b)(x - c)(x - d) \cdots (x - n)}{(a - b)(a - c)(a - d) \cdots (a - n)} \\ & + U_b \frac{(x - a)(x - c)(x - d) \cdots (x - n)}{(b - a)(b - c)(b - d) \cdots (b - n)} \\ & + U_c \frac{(x - a)(x - b)(x - d) \cdots (x - n)}{(c - a)(c - b)(c - d) \cdots (c - n)} \\ & + U_d \frac{(x - a)(x - b)(x - c) \cdots (x - n)}{(d - a)(d - b)(d - c) \cdots (d - n)} + \cdots \quad (12) \end{aligned}$$

A careful examination of the Lagrange formula will reveal the simple regularity in setting up the terms.

Example 5. The solubilities of carbon dioxide in water at various temperatures are given in the table below. Solubility is expressed as grams of carbon dioxide per 100 g of water when the total pressure is 1 atm.

$t^\circ \text{C}$	SOLUBILITY (S)	$t^\circ \text{C}$	SOLUBILITY (S)
0	0.3346	3	0.2978
1	0.3213	5	0.2774

Find the solubility of carbon dioxide at 4° C.

$$S_4 = 0.3346 \frac{(4-1)(4-3)(4-5)}{(0-1)(0-3)(0-5)} + 0.3213 \frac{(4-0)(4-3)(4-5)}{(1-0)(1-3)(1-5)} \\ + 0.2978 \frac{(4-0)(4-1)(4-5)}{(3-0)(3-1)(3-5)} + 0.2774 \frac{(4-0)(4-1)(4-3)}{(5-0)(5-1)(5-3)}$$

$$S_4 = 0.3346 \left(\frac{-3}{-15} \right) + 0.3213 \left(\frac{-4}{8} \right) + 0.2978 \left(\frac{-12}{-12} \right) + 0.2774 \left(\frac{12}{40} \right)$$

$$S_4 = 0.0669 - 0.1607 + 0.2978 + 0.0832$$

$$S_4 = 0.2872$$

The experimental value is 0.2871.

INVERSE NON-LINEAR INTERPOLATION

Inverse interpolation may be performed on non-linear functions by means of either of the Newton formulas or the Lagrange formula. If the Newton formulas are used, proceed in the same way to form the table of differences. For U_x substitute the value of the function for which the corresponding value of x is desired. The result will be a polynomial in x which may be cleared up and solved algebraically if the equation is quadratic, or by one of the approximation methods of Chapter IV if the degree is higher than second. The Lagrange formula may be employed in a similar fashion. Here, however, the resulting equation will be of the $(n-1)$ th degree, where n is the number of terms. Probably the best method of solution will be by successive approximations.

EXTRAPOLATION

The phenomena of nature, and this is predominantly true in chemistry, are very often subject to discontinuities or "breaks" due to change of state or the appearance of new phases. Some examples of discontinuities are: freezing, boiling, separation of solids from solution, and allotropic modifications. As a result, it is very difficult to predict what would happen to a system outside of the range of a variable that has been studied. The attempt to predict, from a set of experimental measurements, the value of a function outside the range for which data are available is known as extrapolation. Mathematically, extrapolation is easily possible, but the interpretation must always be risky. About the best

that can be said for an extrapolation is that it gives the value that would be obtained experimentally *provided no discontinuities occur or no new variables appear*.

Graphical Extrapolation. The graphical method of extrapolation suffices for many purposes where the desired degree of exactness is not very great. A graph of the known values can be prepared on a suitable scale. If there are sufficient points to determine its course with reasonable certainty, the curve can be extended a short distance to obtain extrapolated values of the variable in question. If the curvature of the plotted curve is great the uncertainty in its extension will be correspondingly great. If the function is such that a straight line is obtained, there will be little uncertainty in the extension. However, the original uncertainty in the location of the straight line with relation to the known points will be magnified in the process of extension.

A great many functions can be reduced to a linear form by the use of **substituted variables**. Some of the substitutions most often employed are described in Chapter XI in connection with the general processes of curve fitting. If the function in question can be reduced to a linear form, the extension of the resulting straight line will produce less uncertainty in the extrapolated values.

Extrapolation with Interpolation Formulas. The interpolation formulas given in the first part of this chapter can be used to perform an extrapolation. In this way the computation of an extrapolated value can be carried to as many significant figures as the original data warrant. Extrapolation by this means differs from interpolation only in that the value of x substituted in the formula is a value lying outside the range over which differences are taken.

The great weakness in the use of interpolation formulas involving differences lies in the failure to make use of all the values of the dependent variable that have been tabulated. From an examination of a difference table

U_1			
U_2	ΔU_1	$\Delta^2 U_1$	$\Delta^3 U_1$
U_3	ΔU_2	$\Delta^2 U_2$	$\Delta^3 U_2$
U_4	ΔU_3	$\Delta^2 U_3$	$\Delta^3 U_3$
U_5	ΔU_4	$\Delta^2 U_4$	$\Delta^3 U_4$
.	.	.	.
.	.	.	.
.	.	.	.
.	.	.	.

it is apparent that a difference formula involving U_1 , ΔU_1 , $\Delta^2 U_1$, and $\Delta^3 U_1$ is based on only the first four values of U in the table. An extrapolated value computed from such a formula is therefore a function of only these four values from the table. In general, an extrapolated value computed from a difference formula of r th degree is a function of only $r + 1$ values of the dependent variable.

Other Analytical Methods of Extrapolation. The method of extrapolating with the least numerical uncertainty is one in which values are computed from an equation that represents all the tabulated data. In Chapter XI will be found a discussion of methods employed in fitting an equation to a set of data by the method of averages or of least squares.

In conclusion, we reiterate that graphical and mathematical processes of performing extrapolation are available which involve no great difficulties. Although the process can be carried to almost any degree of numerical exactness, the physical or chemical significance of an extrapolated result must of necessity be doubtful.

PROBLEM SET IX

1. Prove that for a linear function (i.e., first differences are constant) the Newton interpolation formula reduces to the rule of proportional parts.

2. The specific heats of silica glass² at various temperatures are as follows:

$t^\circ \text{C}$	100	200	300	400	500
sp. ht.	0.2372	0.2416	0.2460	0.2504	0.2548

Find the specific heat at 225° and at 350°C .

3. According to Thiesen and Scheel,³ the density of mercury varies with the temperature as follows:

$t^\circ \text{C}$	0	10	20	30	40
$d(\text{g/cc})$	13.5951	13.5704	13.5458	13.5213	13.4969

Calculate the density at 15° and at 22°C .

4. A compilation of values from various handbooks gives the following table for the mobility of the K^+ ion at various temperatures:

$t^\circ \text{C}$	0	25	50	75	100
∞l	40.0	75.0	115	159	206

Calculate ∞l_{18} and ∞l_{40} .

² *Smithsonian Physical Tables*, Eighth Edition, p. 288, 1934.

³ *Tätigkeitiber. Phys.-Techn. Reichsanstalt*, 1897-98.

5. A dish of mercury imbedded in a sand bath was heated by a gas flame. The temperatures observed at 5-minute intervals were:

time (min)	0	5	10	15	20	25
temp. ($t^{\circ}\text{C}$)	20.0	32.5	47.5	65.5	87.0	112.5

What was the temperature at 7 minutes?

6. If, in Problem 5, 89.478 g of mercury were taken for the experiment, from the data of Problems 3 and 5, calculate the change in the volume of the mercury over the first 6-minute interval.

7. The specific gravities of ethyl alcohol-water mixtures at 20°C as determined at the National Bureau of Standards (U.S.A.) are given in the following table. The percentages of alcohol are by weight, and the specific gravities are referred to water at 4°C .

% alcohol	0	2	3	6	8
d_4^{20}	0.99824	0.99453	0.99274	0.98776	0.98470

By divided differences, compute the specific gravity of a mixture containing 5 per cent of alcohol.

8. From the data of Problem 7, calculate by means of the Lagrange formula the specific gravity of a mixture containing 7 per cent of alcohol.

9. At 19.5°C the densities of aqueous solutions of CdCl_2 at various concentrations are given in Hodgman's *Handbook of Chemistry and Physics*. Density (d) is in grams per cubic centimeter and concentration (C) is given as percentage by weight.

C	10	20	30	40	50	60
d	1.087	1.193	1.319	1.469	1.653	1.887

Calculate the concentration of CdCl_2 solution that would have a density of 1.400.

10. In the *International Critical Tables*, Vol. VII, p. 13, the following values are given for the refractive index of water at various wave lengths in the visible spectrum:

$\lambda(\text{\AA})$	4800	4861	5350	5460	5770
n_{20}	1.33750	1.33714	1.33490	1.33447	1.33342

By the Lagrange formula calculate n_{20} at $\lambda = 5338\text{\AA}$. *Hint.* To avoid the necessity of using a six-place log table treat the values of the function as $(n - 1.33)$. There will then be only three significant figures in the function and in the differences.

11. The following table showing the variation of the solubility of SrCl_2 (given as mols of SrCl_2 per 1,000 g of water) with temperature is from the data of Menzies:⁴

$t^{\circ}\text{C}$	145	155	165	175	185	195
S	8.33	8.60	8.89	9.21	9.58	10.03

Solve for the solubility of SrCl_2 at 120°C .

⁴ *J. Am. Chem. Soc.*, **58**, 936 (1936).

12. After comparing the result of Problem 11 with the answer given, consult the original article by Menzies and answer these questions:

- (a) What is the physical significance of the solubility of SrCl_2 at 120°C ?
- (b) What general type of discontinuity does this case represent?
- (c) Make as complete a list as you can of the chemical and physical discontinuities that make extrapolations uncertain.

13. From the tabulated values of the probability integral $\Phi(t)$ given in Table II of the Appendix, calculate to four significant figures the value of t corresponding to $\Phi(t) = 0.50000$. *Important:* Retain this answer for reference in connection with Chapter IX.

CHAPTER VI

THEORY OF MEASUREMENTS

...science forms a closed system simply because it employs the device of cyclic definition.

—J. W. N. SULLIVAN.

STANDARDS

Our entire scientific structure is based upon the presumption that natural phenomena are intelligible. This is a proposition not susceptible of direct proof, but it must be accepted as a doctrinal article. We might restate this as a postulate, *viz.*, natural inorganic¹ systems obey inexorable laws. To discover and establish a natural law, measurements must be made of some characteristic property of a system which may be related to other characteristics of the system in question or to the external conditions of the system. If there is a law relating the given characteristic to other characteristics, it follows that there is a functional relationship which may be expressed as a mathematical proposition. The term "mathematical proposition" must be interpreted here in its broadest aspects in such a manner as Mario Pieri grandly defines mathematics as "the hypothetico-deductive science."² If we limit the argument to the laws of the so-called "exact sciences," which are with few exceptions expressible in mathematical symbols, we can predicate the existence of a mathematical equation (or equations) relating a given property of a system to other of its properties. Granted that a mathematical relation is possible, we can state a corollary to our original postulate (remembering the limitations), *for any measurable physical property there is a unique and true numerical value.*

¹ The word "inorganic" as used here means simply an inanimate system. It must not be inferred that a mixture of organic chemical compounds would not behave in a predictable manner.

² Keyser: *The Pastures of Wonder*, p. 24, Columbia University Press, New York, 1929.

When a scientist has performed a laboratory experiment to measure a physical quantity, he would like to know how good the result is. Comparison of his value with the "true value" would reveal the accuracy of his measurement. Unfortunately neither he nor anyone else knows the "true value"; not even the exact form of the law is known. Accuracy thus defined is a highly elusive attribute. Nevertheless, the words **accurate** and **accuracy** are frequently used in scientific discussion; it should be possible to fix some sensible meanings. Lacking any exact knowledge, we are as usual forced to the best possible approximation, which consists in setting up by definition a series of standards of comparison. Even definitions are subject to human limitations. The task of setting up and defining standards has not been the magic result of the mere resolution to define them. Although a very high state of perfection has been attained, there is room for improvement, just as in every endeavor of man.

Before proceeding further we shall consider briefly the definitions of some standards that have been adopted universally for scientific measurements. Standards may be divided into three classes: fundamental, derived, and relative.

FUNDAMENTAL STANDARDS

The *primary* fundamental standards are those of length, mass, time, and temperature.

Length. The international prototype meter is the unit of length. It is the distance between two fine lines engraved on a certain platinum-iridium bar at the temperature of melting ice. This bar is carefully preserved at the International Bureau of Weights and Measures at Sèvres, near Paris, and highly exact copies have been distributed to the governments of the nations that participated in the Metric Treaty of 1884. The meter was originally intended to represent one ten-millionth of the distance from the equator to the north pole. Subsequent determinations of the earth's meridian quadrant have shown the distance to be more nearly 10,002,100 meters. The distance between the lines on the bar, however, has not been altered. The indestructibility of the meter length has been assured through the determination of the international meter in terms of the wave length of the red line in the cadmium spectrum by Professor A. A. Michelson (1852-

1931). The meter equals 1,553,163.5 wave lengths of this particular radiation.³

Mass. The international kilogram is the standard of mass (weight). It is a platinum-iridium cylinder that was established at the time of the adoption of the international meter, and it is preserved in the same Bureau. The original intention was to have the kilogram equal the mass of a cubic decimeter of pure water at its temperature of maximum density (3.98°C). Later measurements have shown that the water was not free from dissolved air. The correct weight of a cubic decimeter of pure water, freed from dissolved air and weighed *in vacuo* at 3.98°C , is 0.999974 kilogram.

All the units based on both the standards of length and mass are tenfold multiples or submultiples of the standard. Thus the practical unit of length is the centimeter ($\text{cm} = 0.01\text{ m}$), and the practical unit of mass is the gram ($\text{g} = 0.001\text{ kg}$).

Time. The standard of time is the mean solar day, which is the average period of one complete revolution of the earth on its own polar axis. The practical unit of time, the second, is defined

as $\frac{1}{86,400}$ of a mean solar day. The system of units of time has

no simple decimal basis. Rather, it is based on multiples or submultiples of 12 and 60. Since the earliest civilizations, time has been reckoned in this way. The reason lies in the use, by the Babylonians (i.e., Sumerians), of sexagesimal fractions in their astronomical reckonings.

Temperature. Temperature is a manifestation of the random kinetic energy of molecules. By universal convention, temperatures are compared on the basis of the freezing and boiling points of water at a pressure of 760 mm of Hg. In the Centigrade system this temperature interval is divided into 100 units or degrees. The Centigrade scale ($t^{\circ}\text{C}$) is related to the theoretically founded thermodynamic or absolute scale ($T^{\circ}\text{A}$) by the formula $T^{\circ} = t^{\circ} + 273.13$. Inasmuch as any material of which a standard of length can be constructed is subject to variation in dimensions with temperature change, the measurement of lengths, areas, and

³ The decision of the International Conference on Weights and Measures of 1927 was to adopt as the provisional definition of the meter 1,553,164.13 wave lengths of the red light emitted by a cadmium-vapor lamp excited under certain specified conditions.

volumes is seen to be closely related to the exactness with which temperature can be measured and controlled.

The studies of thermometry made by the International Bureau of Weights and Measures showed that mercury-in-glass thermometers were not sufficiently reproducible for very careful measurements. The hydrogen gas thermometer, in which the change of pressure of a constant volume of hydrogen is measured, was adopted as the standard method of measurement.

The universal acceptance of the metric system for scientific measurements marks an epoch in the annals of metrology. Before the Metric Convention, there was a most lamentable diversification of standards for weights and measures throughout the world. In many instances, there was wide disagreement even within a given country. The basis or definition of some of the standards appears almost ludicrous to us now. As examples we may take some of the standards of medieval England. Tradition informs us that Henry I (reigned 1100–1135) with becoming modesty decreed the legal yard as the length of his own arm. Shortly afterward we find demands for more rigorous definition of standards of weights and measures voiced in Magna Charta (1215). As a result we find slight improvement in the time of Henry III. In his famous statute of the Assize of Bread and Ale (1266) we find stated "that by the consent of the whole realm of England, the measure of our Lord the king was made, *viz.*, an English penny called a sterling, round and without any clipping, shall weigh thirty-two wheatcorns in the midst of the ear; and twenty pence do make an ounce, and twelve ounces a pound; and eight pounds do make a gallon of wine, and eight gallons of wine do make a bushel, which is the eighth part of a quarter."⁴ The use of the word "grain" as one of the subsidiary units of modern apothecaries weight arises from this definition. It might also be remarked that the saying "a pint is a pound the world 'round'" has here a foundation since there are eight pints to a gallon. Today, however, the American units, cumbersome as they remain, are legally defined in terms of the metric international standards.

To return to the subject of scientific standards of measurement, let us consider the *secondary* fundamental standards, which are those of area and volume:

⁴ Cf. Hallock and Wade: *Evolution of Weights and Measures*, p. 32, Macmillan, London, 1906.

Area. The standard unit for the measurement of surfaces is the square meter, and the practical unit is the square centimeter (cm^2 or sq cm).

Volume. The standard for the measurement of the volume of a body is the cubic meter or, practically, the cubic centimeter (cm^3 or cc). At the time of the establishment of the kilogram, which was the mass of a cubic decimeter of water, the liter was chosen as the unit of liquid capacity measure. The liter was supposed to be identical to the cubic decimeter. Because of neglecting the dissolved air in the water, the liter, which is defined as "the volume occupied by a kilogram of water at 3.98°C ," is in slight disagreement with the cubic decimeter. Consequently there is a difference between the milliliter ($\frac{1}{1,000}$ of a liter) and the cubic centimeter. The accepted relation between the liter and the cubic centimeter is: 1 liter contains 1,000.027 cc. The liter, however, is derived with the aid of the unit of mass (the kilogram) rather than with the unit of length (the meter); hence its subdivisions are, strictly speaking, called milliliters (ml) instead of cubic centimeters (cc). The error in assuming the cubic centimeter equivalent to the milliliter is only 1 part in 37,000, and hence it has no practical significance in analytical chemistry. But, for the sake of consistency, the liter and the milliliter are used as the units of volumetric measurements. Many of the recent textbooks of chemistry have adopted this nomenclature, but many are still in circulation which continue to confuse cubic centimeter with milliliter. It is to be hoped that revisions will correct this inconsistency.

DERIVED STANDARDS

From these fundamental standards we have just defined, the second great class, derived standards, is deduced. To list and discuss all of them would be outside the scope of this book. We shall consider only a few that are used most frequently by chemists, namely: density, specific gravity, energy, pressure, and heat quantity.

Density. The absolute density is the weight in grams of a cubic centimeter of material at a fixed temperature. The absolute density of water thus defined is 0.999974 at 3.98°C . If density were defined as grams per milliliter, water would have an absolute

density of 1.000000... at the temperature of maximum density. Density is denoted $d_{t^{\circ}}$.

Specific gravity is defined as the ratio of the density of a substance at a specified temperature (frequently 15° C) to the density of water at 4° C. Thus water at 4° C has a specific gravity of unity. Specific gravity is denoted by $d_4^{t^{\circ}}$.

Energy. The unit of energy (work) is defined as a dyne acting through a net displacement of a centimeter. The dyne is specified as the force required to change the velocity of a mass of a gram by one centimeter per second in one second. This standard unit of energy or work is called the erg.

Pressure is force per unit area. It is defined as dynes per square centimeter. The practical unit of pressure is the atmosphere (atm), and this is taken to represent the average pressure of the atmosphere at mean sea level at 45° latitude. The atmosphere by international agreement is specified as the pressure required to support a column of mercury 760 mm in height in an evacuated tube at 0° C.

Heat Quantity. The unit of heat energy is the calorie (cal). It is the quantity of heat required to raise the temperature of one gram of water from 3.5° to 4.5° C. Another unit known as the mean calorie has been defined as $\frac{1}{100}$ of the quantity of heat required to raise the temperature of one gram of water from 0° to 100° C. The calorie and the mean calorie are slightly different because the heat capacity of water varies with the temperature. The large calorie or kilogram calorie (Cal) is the quantity of heat required to raise 1,000 grams of water from 3.5° to 4.5° C.

RELATIVE STANDARDS

Some physical quantities are not susceptible of measurement in terms of any absolute standards, at least not with the techniques that have been developed up to the present. It is usually possible, however, to obtain a numerical relation or ratio between the values of the given property for two substances *A* and *B*. Then if *B* can be related to *C*, *C* is related to *A*. Obviously there is no limit to the extension of such a system of relations or ratios except the number of substances available.

This situation might be illustrated by a mathematical analogy. Given two linear equations containing three unknowns,

$$a_1x + b_1y + c_1z = d_1$$

$$a_2x + b_2y + c_2z = d_2$$

there is no possibility of a complete solution for all three variables. There are two courses open: (1) solve for each of two unknowns in terms of the third variable; (2) assign an arbitrary value to one unknown and solve for the other two. In scientific practice the second of the two courses is adopted. An arbitrary value for the property in question is assigned to one substance designated as the standard or to one particular state of all substances designated as the standard state.

Atomic Weights. The most familiar example of an arbitrarily assigned relative standard in chemical science is the assignment of 16.0000 as the atomic weight of oxygen. All other atomic weights are based upon oxygen as the standard substance with its fixed weight. With the realization that molecular weights are simply sums of atomic weights and that the mol is the chemical unit of weight, it is possible to appreciate the far-reaching significance of the choice of this standard.

Electrode Potentials. A metal in contact with its ions in solution has an electrical pressure or potential. This potential exists only as a tendency, unless two such systems are placed in electrical contact. Then a resultant potential of the two systems can be measured which expresses the difference between the electrical tendencies of the two systems, or half-cells, as they are called. It is important to know the values of these electrical tendencies since they are a measure of chemical affinity or reactivity. Here the choice has been made to designate the hydrogen electrode as having a half-cell potential of zero. Thus any electrode system may be measured against the hydrogen electrode, and the potential observed will be the potential of the given electrode system. The hydrogen electrode is specified as hydrogen gas at one atmosphere pressure bubbling over a platinized platinum surface immersed in an aqueous solution containing unit concentration of hydrogen ions. For complete rigor, unit *activity* of the hydrogen ions must be specified. Activity is the idealized or thermodynamic concentration.

Activity. It happens that activity cannot be related to absolute standards, and the necessity of a relative standard is imposed. This is accomplished by the designation of standard states. A pure solid at unit pressure is specified as having unit activity. A solute in an infinitely dilute solution is designated as having activity equal to the concentration; in other words, the activity coefficient is unity. A gas has unit activity when it is in the condition of unit fugacity. Fugacity is the thermodynamic pressure or the pressure the gas would have if it were ideal. Fugacity, in turn, requires a relative standard for definition. From this it is apparent that the scientific structure is a highly interrelated and interdependent set of definitions and standards. The situation is reminiscent of a famous verse that has been quoted often and undoubtedly will be quoted many more times. That verse is:

So Naturalists observe, a flea
Has smaller fleas that on him prey;
And these have smaller still to bite 'em,
And so proceed *ad infinitum*.⁵

MEASUREMENT

In nature nothing is ever entirely independent. Every phenomenon is the result of a multitude of contributing causes which we may call variables. The behavior of the smallest atom is the resultant of its being acted upon by every other atom in the universe. We must not overlook, however, that this particular atom possesses certain atomistic characteristics that determine the way in which it will be affected by all the other atoms. Substitute the word "planet" for "atom" in the preceding statements, and they remain equally true. Thus, one atom cannot be singled out and its properties ascribed to itself and itself alone. Any such statement must be at least an implicit expression of its environment at a particular time. The study of atoms is not so complicated as a

⁵ Swift, "On Poetry." Found in *The Works of Jonathan Swift*, Second Edition, p. 311, Vol. XIV, Bickers and Son, London, 1883. Sometimes confused with a verse by Augustus de Morgan, *A Budget of Paradoxes*, p. 377, Longmans, Green and Co., London, 1872, which goes:

Great fleas have little fleas upon their backs to bite 'em,
And little fleas have lesser fleas, and so *ad infinitum*.
And the great fleas themselves in turn, have greater fleas to go on;
While these again have greater still, and greater still, and so on.

first glance might indicate. To all intents and purposes only the immediate neighbors of a particular atom exert an appreciable effect. All forces are known to diminish with distance, and the effects of far-removed neighbors become vanishingly small. This is equivalent to saying that a portion of the universe may be set aside and its properties ascribed to a few preponderating variables. As an example, consider astronomical science. It is known that electrostatic forces exist among the planets of our solar system, but the electrostatic forces are so small in comparison with gravitational forces that they are negligible. Hence the calculations of orbits, eclipses, etc., are based on gravitation, neglecting electrostatics.

If we set out to observe a particular phenomenon and make quantitative measurements, the first thing to be done is to find out which of the variables that condition the phenomenon are ponderable. The fact that the reverse of this procedure is the one usually followed in practice does not detract from the validity of the argument. The second step consists in devising surroundings such that most of the variables will be fixed or, if not fixed, that provision may be made to apply corrections that serve the same end. This fixing of variables is necessary because, as a rule, scientific techniques are such that only two variables may be studied simultaneously. Then, going further, one of these two variables is "fixed" at certain "known" values and corresponding values of the other variable are observed. From the data thus obtained, a law relating the two variables may be deduced. Successively, taking two variables at a time, a series of "partial" laws may be established so that every ponderable variable may be related to each of the other ponderable variables. If all these "partial" laws can be combined to give a "composite" law, we have then obtained a typical scientific law. Because the variables which we have decreed inappreciable are not taken into account, the law is still incomplete. This is what was meant by the statement that occurs earlier in this chapter, "the exact form of the [natural] law is not known."

An example of this combining of "partial" laws that is familiar to every student of elementary chemistry is the combination of the laws of Boyle and of Charles to yield the ideal gas law, $pV = RT$. This law is incomplete because it considers only three variables. The gas law of van der Waals is a better representation because it

takes into account the variation of attractive forces between molecules and the volume of the molecules, but even this law,

$$\left(p + \frac{a}{V^2}\right)(V - b) = RT,$$

is far from perfect.

The production of a law is not the result of one experiment; rather it is an evolutionary process. The first stage is the advancement of an hypothesis as a basis for reasoning and experimentation or as an explanation of certain observed facts. If supporting evidence from other sources can be obtained, the hypothesis may progress to the status of a theory. Finally, if an exhaustively critical examination of a theory shows that it is valid for all cases that can be reasonably investigated, that theory is accepted as a law.

OBSERVATIONS

First let us consider an important experimental fact. *Successive observations of a physical quantity disagree*, provided that each observation is carried out as exactly as the method allows so that the values are not rounded off. A carpenter in measuring a door-step with a foot-rule would measure from left to right and report, say, 3 ft 2 in. If asked to check his measurement he would probably measure from right to left and report 3 ft 2 in. again, neglecting any fractions of an inch. Although his measurement would be sufficient for the purpose of computing the board feet of lumber it is still a rounded value. Suppose that a coppersmith were measuring that same door-step in order to fit it with a sheet copper tread; his measurement would be more refined. He would measure, say, 3 ft 2 $\frac{3}{64}$ in. If, however, he were to measure again, or his assistant checked him, the result might be 3 ft 2 $\frac{3}{64}$ in., or perhaps the fraction might be $\frac{4}{64}$ in.

In a chemical laboratory, successive weighings of a platinum crucible under supposedly identical conditions gave the following weights in grams:

19.5437
19.5434
19.5436
19.5438

To what are these variations due? In the first place, the scale

over which the pointer swings is graduated only to a certain degree of fineness, and estimations between the scale marks are necessary. Anyone with a little practice can learn to estimate the scale divisions within one- or two-tenths of a division. This uncertainty might account for a variation of ± 1 in the fourth decimal place of the result. The variation is of the order of several tenths; how then can we account for it? The answer lies in the neglect of further variables deemed to exert a negligible effect and also in the failure to fix completely the ponderable variables. In theory, the act of weighing with the analytical balance is a study of the action of a delicately poised compound pendulum. We load the two sides of the pendulum as nearly equally as possible (the addition and removal of weights is the independent variable) and observe the deflection of the pendulum from its equilibrium position of rest. The deflection observed is the dependent variable. Meanwhile we have attempted to fix all the other variables, and, assuming that the ordinary rules of good practice in weighing have been observed, these omissions have been allowed: (1) incomplete control of the quantity of moisture adsorbed on the crucible and weights; (2) minute but appreciable temperature variations inside the balance case; (3) the possibility of small electrostatic charges on the crucible or the weights; (4) failure to take into account the variable quantity of dust particles in the air which might settle on the crucible or weights at different rates from time to time; and (5) using different combinations of weights in different weighings, for, even though the weights have been calibrated, the calibration cannot be perfect.

The free operation of these five variables, and there may be others not yet recognized, contributes to the formation of accumulated discrepancies which may be large or small or positive or negative deviations from the "true" value. The laws of chance would predict that, with a group of practically independent variables freely operating, more times than not they would tend to cancel one another and hence produce only small deviations. The combination of circumstances necessary to produce a large deviation would be unlikely; hence, small deviations are more likely to occur than large ones. The laws of chance would further predict that positive and negative deviations would be equally likely, since the variables are operating freely.

The foregoing paragraph is essentially a statement of the

"theory of errors." This subject will be considered in detail in Chapter IX.

AVERAGES

Assuming the free operation of all the further variables and also the predictions concerning the sizes and signs of the deviations just stated, it follows that the greater the number of observations made the greater the chance of obtaining an exactly equal number of plus and minus deviations of like sizes that would cancel each other if all the observations were averaged together. We are now in a position to dispose of the bothersome and elusive term, "true value." We shall define the "true value" as the average of an infinite number of observations. To return to practical cases, we shall assume that, for a finite number of observations, the average gives the "best value," in other words, the best approximation to the "true value."

ERRORS

Strictly speaking, the only correct definition of error is the discrepancy between an observed value and the "true value." We never have a "true value"—only an average of a finite number of observations. We shall define the difference between a single observed value and the average as the **deviation** of that observation. Owing to another of those unfortunate inconsistencies which have crept in during the evolution of our scientific vocabulary the terms "error" and "deviation" have become almost hopelessly confused. The terms are used interchangeably in scientific literature, and an attempt to rectify the mistake at this late date could result only in further confusion.

Through the expression of the deviation of an observation we obtain a measure of the "goodness" of that observation. We could indicate the goodness of a series of observations by stating the value of each observation and its deviation. For ready reference to a large number of observations such a scheme would be impractical and we shall need a method of expressing the average deviation in a series of results. This average deviation tells how consistent the results are among themselves. Mathematical expressions for dispersion measures will be treated in Chapter VIII.

PRECISION

The difference between accuracy and precision can be illustrated by a series of targets or "bull's-eye" diagrams. In Fig. 7 we see three such targets. Rifleman *A* shows the ideal combination of high precision and good accuracy. Rifleman *B* displays high precision but poor accuracy. His shots are well bunched, but, because of incorrect sight adjustment or an unreckoned wind velocity, the center of impact of his shots is well away from the center of the target. Rifleman *C* displays poor precision. His accuracy may or may not be good; it is probably an accident

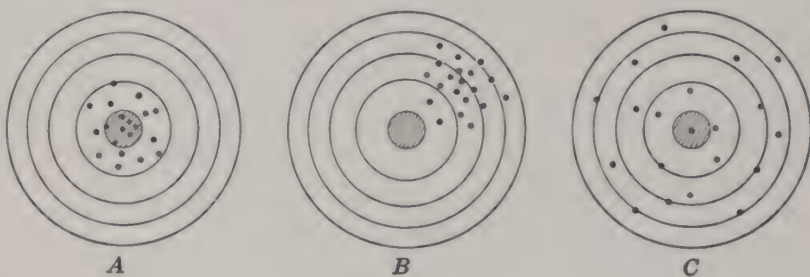


FIG. 7.

if the center of impact of his shots agrees with the center of the target. He is an inconsistent marksman, and we should rate his ability very low. It might have been worse if he had been blindfolded! We conclude that *A* and *B* are consistent marksmen but that *B*'s shots are grouped away from the bull's-eye because of a **constant error**. The best-planned series of measurements may be subject to the same phenomenon because one ponderable source of variation has been overlooked. It is like the one slip in a "perfect crime." It is because of the possibility of constant error that we must ever be on guard against accepting a "best value" (average) as correct. Rifleman *B* has the record of his 20 shots on the target, and he can locate the source of error (adjustment of sights, wind velocity, etc.). In scientific measurements, we have no bull's-eye other than the imaginary "true value" with which to compare the "grouping of shots." Scientific measurement, to carry the analogy further, is an attempt to locate a bull's-eye by shooting at a blank target with a carefully adjusted gun. At least, the experimenter always hopes the gun is properly adjusted.

ACCURACY

In conclusion, let us qualify the words "accurate" and "accuracy." When the adjective "accurate" is employed to describe the measurement of a physical quantity, it must be qualified with a modifying term such as "highly accurate" or "fairly accurate" because the word alone would signify "completely accurate." Likewise with "accuracy" we must specify a degree such as "high (or low) degree of accuracy." The only conceivable sense in which the phrase "absolute accuracy" can be used is in respect to defined standards. The statements, "The atomic weight of oxygen is 16.0000" or "One meter equals 100 centimeters" are among the few examples of absolute accuracy.

PROBLEM SET X

1. The gram-molecular volume is approximately 22.4 l at 0° C and 760 mm pressure. Given the following relations:⁶

$$1 \text{ oz avoirdupois} = 28.350 \text{ g}$$

$$1 \text{ cu ft} = 28.317 \text{ dm}^3$$

Compute the ounce-molecular volume in cubic feet.

2. According to the legal definition of the U.S.A. inch, there are exactly 39.37 in. to the meter. Find the number of square feet in a square meter.

3. The Centigrade scale of temperature (°C) is related to the Fahrenheit scale (°F) by the expression:

$$^{\circ}\text{C} = \frac{5}{9} (^{\circ}\text{F} - 32)$$

If the absolute zero of temperature is -273.13°C , what is it on the Fahrenheit scale?

4. Convert -40°F into degrees Centigrade.

5. The British thermal unit (Btu) is defined as the quantity of heat required to raise the temperature of one pound of water at its maximum density through one degree Fahrenheit. Recalling that there are 16 ounces to a pound, use the information in Problems 1 and 3 to compute the factor for converting Btu to calories.

6. At 15°C the absolute density of water is 0.99910 g per cc. What is the absolute density of water at that temperature expressed in pounds per cubic foot?

⁶ This approximate coincidence was first noted by J. W. Richards, late Professor of Metallurgy, Lehigh University.

CHAPTER VII

CLASSIFICATION OF ERRORS

Now, that which imparts truth to the known and the power of knowing to the knower is what I would have you term the idea of the good, and this you will deem to be the cause of science, and of truth in so far as the latter becomes the subject of knowledge; . . .

—PLATO.

We have already seen that the word "error" suffers the burden of several meanings. It covers, in present usage, not only "error" as rigidly defined and "deviation" which is an indication of the precision but also "discrepancies" which are due to faulty technique in experimentation. It is possible for the deviation of an observation to be the error of that observation, if the value from which the deviation occurs is correct by the definition of a standard. Otherwise, the deviations may approach the errors according as the number of observations is increased.

Errors may be roughly classified as:

1. Those about which we may do something.
2. Those about which we can do nothing.

The first class we shall designate as **corrigible errors**, which means it is possible to correct for them. The second class we shall designate as **random errors**. Several writers subdivide errors of the corrigible kind into "constant" and "systematic." The distinction is more nearly one of degree than of kind, and, since "constant" errors are not always constant, we shall reserve the term for a specific purpose—as the interpretation of Rifleman *B*'s results in the preceding chapter (p. 94).

In the remaining chapters of this book we shall be concerned with a study of statistics and the laws of chance, the principles of which we invoke in an attempt to interpret the degree of "accuracy" of a measurement in terms of its random errors. Errors of the corrigible kind are not subject to the laws of chance, and what

we "do about them" is to attempt to eliminate them through theoretical considerations and by repeating measurements under a great variety of conditions, using as many different techniques as possible.

CORRIGIBLE ERRORS

To familiarize the reader with the idea of corrections and how they are applied, we shall discuss in some detail a few of the types of measurement most often performed in the chemical laboratory. These may serve as a model in planning a series of measurements or in criticizing the work of another observer. Obviously, it would be impossible to discuss each of these types in sufficiently complete detail for the discussion to serve as directions for the instruction of a beginner. A book much larger than this one could be written about any one of the succeeding types of measurement. It must be presumed, therefore, that the reader is familiar with the basic points of acceptable laboratory technique as given in an experimental manual. The discussions that follow are concerned only with refinements in processes of measurement. In other words, the sources of gross error attributable to ignorance or awkwardness will be passed over.

TEMPERATURE

The fundamental basis of temperature measurement has already been defined as the kinetic energy difference in water at its freezing and boiling points, respectively (p. 84). It has also been pointed out that the standard method for measuring temperatures is the gas thermometer. For most practical purposes the gas thermometer is too bulky and complicated for ready use, whereas the mercury-in-glass type of thermometer, though much less exact, offers the advantages of simplicity of operation and small bulk. Fortunately it is possible by means of simple corrections to rectify to a high degree the lack of exactness of the mercurial thermometer and thereby bring results into close agreement with gas-thermometer measurements.

In the first place, in temperature measurements, by whatever method, it is preferable to arrange the system of which the temperature is to be taken so that its immediate surroundings will be at approximately the same temperature. This is conveniently accomplished by means of a heating or cooling bath which, for

highest accuracy, should be an automatically regulated thermostat. Such a bath prevents sudden fluctuations in heat radiation and convection from or into the system. If a bath is not practicable, the system may be surrounded by an insulating material the heat conduction of which is very low. Such an arrangement is sometimes just as satisfactory as a bath.

The inexactness of the mercurial thermometer arises from several sources, namely: (1) incorrect location of engraved scale and non-uniform bore of tube; (2) difference between the expansion of the glass and the expansion of mercury; (3) failure to expose entire thermometer to system being measured so that a temperature gradient exists along the stem; and (4) time lag in contraction of glass following expansion. We shall consider the remedies for these faults in the order named.

Calibration. A manufacturer, no matter how conscientious, can not produce thermometers that are entirely uniform. The very nature of glass working prohibits that. Without trial, no thermometer can be accepted as being capable of registering correctly the so-called fixed points of thermometry. The process of locating these points in terms of the scale engraved on the thermometer is known as calibration. The most commonly used fixed points are the freezing and boiling points of highly purified water. This is appropriate since they are the fundamental standards. If the thermometer is submerged to the zero mark in a properly prepared mixture of crushed ice and water and allowed to reach equilibrium, the reading observed will give the zero point of the thermometer. Here it is a simple matter to expose the entire mercury column to the ice-water system. Of course, several trials should be made to insure equilibrium. To locate the boiling point of water on the scale, the value of the boiling point for the prevailing barometric pressure must first be determined by means of the formula:

$$t = 100 + 0.0375 (b - 760) \quad (1)$$

in which b is the observed pressure in millimeters of mercury. Mere comparison of the temperature recorded with the boiling point computed by equation (1) gives the second fixed point, provided that the entire length of the mercury thread is submerged in the boiling liquid or the vapor, depending on the method. If part of the stem is not so exposed, an emergent stem correction

must be applied to the observed temperature before comparison with the calculated value is permissible.

If the tube had uniform bore, the scale corrections for all temperatures between the freezing and boiling points could be obtained from a linear graph. To prove the uniformity it is necessary to check the intermediate range with other fixed points, such as the transition point of sodium sulfate decahydrate, $\text{Na}_2\text{SO}_4 \cdot 10\text{H}_2\text{O}$, to anhydrous sodium sulfate, Na_2SO_4 , at 32.384°C ; or the boiling point of ethyl alcohol at 78.26°C to which must be applied the barometric correction

$$t = 78.26 + 0.034 (b - 760) \quad (2)$$

Other fixed points below 0°C and above 100°C may be used to cover the whole range of the thermometer in question.¹ With a sufficiently large number of such points correctly determined, a table of corrections for the given thermometer may be prepared covering the range for which it is to be used.

Another method, although less desirable for many purposes, consists in the comparison of a given thermometer with one which has been calibrated by a highly skilled worker as, for example, a thermometer with a National Bureau of Standards (U. S.) certificate of calibration. The comparison may be accomplished by immersing the two thermometers close together in a suitable liquid bath equipped with a stirring device and obtaining a direct comparison at various temperatures. Of course, it will be necessary to apply the emergent stem correction to each thermometer.

Emergent Stem Correction. By emergent stem is meant that portion of the mercury thread which is not exposed to the system the temperature of which is being measured. The general form of the formula for emergent stem correction is given by

$$ka (t - t_0) \quad (3)$$

¹ A long list of fixed points is given in a table entitled "Fixed Temperatures for Thermometer Calibration" in Hodgman: *Handbook of Chemistry and Physics*, Chemical Rubber Publishing Co., Cleveland, revised annually.

The National Bureau of Standards has prepared for distribution several substances to serve as thermometric standards. The melting points of these substances have been carefully determined, and the Bureau furnishes a certificate with each sample. See *Supplement to the National Bureau of Standards Circular C398*, which is sent free upon request.

where k is the so-called cubical coefficient of expansion of mercury-in-glass and is equal to the difference between the coefficients for mercury and glass. The average value of k is usually given as 0.00016.² In the formula, a represents the number of thermometer scale divisions of the mercury thread not exposed to the system, t is the observed temperature of the system being measured, and t_0 is the temperature of the emergent stem. Actually a temperature gradient will exist along the emergent stem. As a good approximation, the mean temperature is taken by suspending an auxiliary thermometer with its bulb touching the recording thermometer at the middle of the emergent stem. In final form the formula for temperature corrected for emergent stem is

$$t_{(\text{corr.})} = t + 0.00016a (t - t_0) \quad (4)$$

Glass Fatigue. Glass is a supercooled liquid, not a crystalline solid, hence we would expect anomalistic behavior. Upon heating, glass expands in a fairly regular fashion, but upon cooling the deformation remains, in part at least, and oftentimes months are required for the return to normal shape and size. Therefore a thermometer which has been used for high temperatures may be in error by as much as 1° C when it is used at lower temperatures, because of the "fatigue" of the glass bulb. If wide ranges of temperature are to be observed, it is best to have several thermometers, one for each portion of the range of temperatures. If such a choice is made, each thermometer should be "baked" for several hours at a middle value of its range before calibration and again each time before use.

Electrical Measurement of Temperature. The mercurial thermometer can be used over only a limited range of temperature. Mercury freezes at -38.87° C and boils at 356.9° C. By filling the stem with nitrogen gas under high pressure the upper limit may be extended well above 360°. If quartz is substituted for glass, the range may even be extended to 650° C. For very high and very low temperatures the most commonly used device for measurement is the thermoelectric couple. The exactness of the measurement is limited chiefly by the sensitivity of the apparatus for measuring the emf developed. For the careful measurement of small temperature differences a thermopile, consisting of an appro-

² For values of k for different glasses at various temperatures, see *International Critical Tables*, Vol. I, p. 56.

appropriate number of thermocouples in series, is often employed. Usually, however, the chemist chooses the Beckmann thermometer for small temperature differences. This thermometer can be read to 0.001°C .

For extended treatments of thermometric calibration and corrections the reader may consult Rowland, "On the Mechanical Equivalent of Heat," *Proceedings of the American Academy of Arts and Sciences*, **15**, 75-200 (1880); or Wood and Cork: *Pyrometry*, McGraw-Hill, New York, 1927; or Burgess and Le Châtelier: *Measurement of High Temperatures*, Third Edition, Wiley, New York, 1912.

PRESSURE

The unit of pressure, the atmosphere, has been defined as the pressure required to support a column of mercury 760 mm in height in an evacuated glass tube. The further conditions of the definition specify a temperature of 0°C and that the measurement be made at sea level and at 45° latitude where the gravity constant has the value $980.665\text{ cm per sec}^2$. When pressures are measured by means of the mercurial barometer under conditions other than the above, several corrections must be applied to give correct results. The errors due to adjustment of the level of the mercury reservoir and to parallax in reading the height of the column need no detailed discussion, since compensation for them is provided in the construction of good barometers.

Temperature Correction of Barometer Height. When the height of a barometer is observed at a temperature other than 0°C , the observation will be in error because of the expansion of both the mercury and the brass scale on which the height is read. The mean coefficient of cubical expansion of mercury is 0.0001818 per degree. If t is the temperature of the barometer and b is the observed height, the pressure (p) corrected to 0°C for mercury expansion is given by

$$p = b - 0.0001818bt \quad (5)$$

The linear coefficient of expansion for brass is 0.0000184 per degree, and since the brass scale expansion would give a low value, the correction

$$0.0000184bt$$

must be added to the observed height b . The combined correction for mercury and brass expansion³ is given by

$$p = b - (0.0001818 - 0.0000184)bt \quad (6)$$

At temperatures well above 0° C, the vapor pressure of mercury⁴ in the barometer tube is sufficient to produce a slight pressure which causes a corresponding depression of the mercury column. To correct for this depression it is necessary to add to the observed height (in millimeters) the vapor pressure of mercury (in millimeters). It is only above 50° C that this correction is appreciable in barometry or manometry as usually practiced.

Correction for Capillary Depression. Since mercury does not "wet" glass, the surface tension results in a convex meniscus that is depressed, in a tube small enough for capillarity to display an appreciable effect. If the diameter of the barometer tube is 25 mm or larger the depression will not be appreciable. For tubes smaller than 25 mm, in order to find the correction, it is necessary to know the internal diameter of the tube and the height of the meniscus itself. The height of the meniscus must not be confused with the height of the mercury column. For example, it can be ascertained from the table of depressions⁵ that, for a barometric tube 8 mm in diameter in which the mercury meniscus is 0.6 mm high, the correction which must be added to the observed barometer height is 0.35 mm of Hg.

Correction for Gravity. It is only at a latitude of 45° and mean sea level that the gravity constant equals 980.665 cm per sec². At other latitudes and altitudes, the barometric height does not give a correct estimate of the pressure. This is true because of variation in the weight of the mercury column, although, of course, its mass is invariant. The correction factor is given by

$$1 - 0.0026 \cos 2\phi - 0.0000002H \quad (7)$$

in which ϕ represents the latitude and H is the altitude expressed in meters above mean sea level. The correction factor gives the value of the ratio of g at ϕ and H to g under standard conditions.

³ Tabulated values of the corrections that are to be subtracted from observed heights are given in *I.C.T.*, Vol. I, p. 69, Table 1B.

⁴ See *I.C.T.*, Vol. III, p. 206.

⁵ See *I.C.T.*, Vol. I, p. 70, Table 2.

The barometric height must be multiplied by the value of the correction factor to give the corrected pressure.

WEIGHING

The process of weighing with the analytical balance has already been discussed in Chapter VI but not with respect to the corrigible errors of weighing. The analytical balance gives a determination of the mass of an object. Since the process consists of direct comparison, this method of weighing is independent of latitude and altitude. Assuming that a given balance is in first-rate working order, that it has been properly adjusted, and that the operator is familiar with the principles of acceptable technique, the sources of corrigible error in weighings with the analytical balance are: (1) imperfect adjustment of, and inconstancy of the set of weights, (2) inequality of the balance arm lengths, and (3) buoyant effect of the air when weighed object and weights have different densities. The methods of correcting for these errors will be considered in the order in which they are listed.

Calibration of Weights. The production of a set of weights, each of which is so nearly correct that no error could be detected by means of the analytical balance, would involve an almost prohibitive price. Furthermore, assuming such a set of weights to have been made, the only way in which the set could be maintained in this state of perfection would be to seal it in a vacuum or in a non-corrosive atmosphere and never use the weights. If all our weights were treated in this fashion, we should be reduced to the state of medieval science and chemists would then sit in easy chairs and speculate about the "heaviness" of objects.

To return to practicality, the weights available for use are not exactly what their face values represent them to be. Several methods of weight adjustment or calibration are in use in various laboratories. One method consists of a relative calibration in which, by either a double substitution procedure⁶ or the Richards' procedure,⁷ n equations are obtained relating each of the n weights of the set to one weight (usually the 10-g) which is assumed to be correct. The solution of the n equations gives a correction for

⁶ P. F. Weatherill, *J. Am. Chem. Soc.*, **52**, 1938 (1930).

⁷ T. W. Richards, *ibid.*, **22**, 144 (1900); Hopkins, Zinn, and Rogers, *ibid.*, **42**, 2528 (1920).

each of the weights which may be $+$ or $-$ depending on whether the weight is heavier or lighter than its face value. The set of weights when used with the application of the corrections is consistent within itself and gives weighings which are *relatively* correct. If a set of primary standard weights (see below) is available, the weight which was chosen as the relative standard can be compared with a weight from the primary standard set and its "absolute" weight determined. Upon the solution of the n equations relating each of the n weights to the absolute value of the 10-g weight, absolute corrections for the set will be obtained.

Another method of calibration consists in procuring a set of primary standard weights. Usually (in the United States) a carefully manufactured set of high-grade weights, made of bronze or brass and plated with gold or platinum, is purchased and sent to the National Bureau of Standards for comparison with the sub-multiples of the national prototype standard. The table of corrections thus obtained renders the set "absolutely" correct. The weights to be calibrated can be compared directly with the corresponding primary standard weights, and their absolute corrections can be determined. In most laboratories, the primary standard set is kept in a protected place and compared at intervals with a secondary standard set. The secondary set then serves as a working laboratory standard. The intervals between comparisons will be dictated by the degree of exactness desired and the durability of the secondary set.

In any set of weights the durability determines the length of time between subsequent calibrations. The most constant weights are usually those made in a single piece and plated with a corrosion-resistant noble metal. Single-piece weights made of fused quartz have been employed as primary standard weights. But an objection to them is the likelihood of the acquisition of an electrostatic charge when a quartz weight is pulled from a felt-lined weight box. The majority of weights are made with a separate knob which screws in and out to give access to the hollow center in which is placed fine shot or metal foil to give the final adjustment in manufacture. Nearly always, the shot or foil is subject to slow oxidation by atmospheric oxygen which leaks in, and the gain in weight thus produced is one of the greatest sources of inconstancy. Weights may be lacquered to prevent corrosion, but the lacquer takes up and loses moisture with changes in the humidity.

Correction for Inequality of Balance Arms. Unless the distances between the center knife-edge and the knife-edges at the ends of the balance beam are the same, the correct mass of an object will not be given by the sum of the weights that balance the object. A correction factor to eliminate the effect of balance-arm inequality can be obtained by the method of double weighing. An object (a weight from another set may be used) is weighed first by placing it on the right-hand pan and counterbalancing it with weights. Then, placing the object on the left-hand pan, the procedure is repeated. Denoting w_R and w_L as the apparent weights of the object in the right and left pans, respectively, R and L as the lengths of the respective arms, and W as the real weight of the object, then, by the law of the lever,

$$(1) \quad WL = w_LR$$

$$(2) \quad w_R L = WR$$

Multiplying (1) by (2)

$$Ww_R L^2 = Ww_L R^2$$

Or

$$\frac{R^2}{L^2} = \frac{w_R}{w_L}$$

The right side of the equation will be unchanged if we add 1 and subtract w_L/w_L .

$$\frac{R^2}{L^2} = 1 + \frac{w_R}{w_L} - \frac{w_L}{w_L}$$

whence,

$$\frac{R}{L} = \sqrt{1 + \frac{w_R - w_L}{w_L}}$$

A familiar approximation is

$$(1 + a)^{1/2} = 1 + \frac{1}{2}a - \frac{1}{8}a^2 + \frac{1}{16}a^3 - \dots$$

If a is very small, all terms beyond the second are negligible. Hence

$$\frac{R}{L} = 1 + \frac{w_R - w_L}{2w_L} \quad (8)$$

It is customary to weigh by placing the object to be weighed on the left pan and to counterbalance by weights on the right pan. When the weighing has been so conducted, the corrected weight of

the object is obtained by multiplying the sum of the weights by the value of the ratio R/L . This conclusion follows from the first equation of the lever

$$WL = w_L R$$

because

$$W = w_L \frac{R}{L}$$

Unfortunately this ratio for a given balance does not remain completely constant, and the ratio is slightly different for different loads. For most purposes, a frequent redetermination of the ratio at various loads will suffice. For extreme exactness, the double weighing method of Gauss⁸ is usually used.

Correction for Buoyancy of the Air. The famous principle, discovered by Archimedes while he was bathing, stated in the words of today is: a solid body immersed in a fluid is subject to a lifting force which equals the weight of the fluid displaced. A vacuum is a fluidless medium; hence an object will weigh less in air than it would in a vacuum by the amount of the buoyant force due to air displacement. In weighing, if the weights and the object being weighed displace the same weight of air (that is, the two have the same density), then no error is introduced. When the densities are not the same, then the buoyant forces will be different and the result will be in error to the extent of the difference between the two buoyant forces. Denoting the apparent weight of the object and of the weights balancing the object as W and the densities of the object and weights as d_1 and d_2 , respectively, for a good approximation the volumes are

$$V_{\text{object}} = \frac{W}{d_1}; \quad V_{\text{weights}} = \frac{W}{d_2}$$

Denoting the weight of 1 cc of air by a , the weights of air displaced are:

OBJECT	WEIGHTS
wt. of air = aV_{object}	wt. of air = aV_{weights}
$= a \frac{W}{d_1}$	$= a \frac{W}{d_2}$

⁸ Gauss's method is fully described in Fales and Kenny: *Inorganic Quantitative Analysis*, New Edition, p. 101, Appleton-Century, New York, 1939, and in Kolthoff and Sandell: *Quantitative Inorganic Analysis*, p. 203, Macmillan, New York, 1936.

The buoyant force on the object is the difference between the two buoyant forces,

$$aW \left(\frac{1}{d_1} - \frac{1}{d_2} \right) \quad (9)$$

This correction term must be added to the apparent weight. An average value⁹ of a that suffices for most purposes is 0.0012 g.

Hence,

$$W_{\text{corr.}} = W + 0.0012 W \left(\frac{1}{d_1} - \frac{1}{d_2} \right) \quad (10)$$

Equation (10) gives the form for the correction of weights to *in vacuo*.

Gravimetric Analysis. The foregoing discussion of weighings has neglected the possibility of changes in the object being weighed, either during the process of weighing or between successive weighings. In chemical laboratories most of the weighings performed are those involved in gravimetric analyses where the possible variations in the weights of objects being weighed constitute even greater sources of error than those already listed. In gravimetric analyses, a single weighing means nothing! Successive weighings alternating with successive dryings or ignitions must be made until a constant weight is obtained.

The failure of a crucible (or other vessel) and its contents to weigh the same on successive weighings may be due to one of the following causes:

1. **Incomplete volatilization of a volatile product.** In many gravimetric determinations, it is necessary to decompose the precipitated product to produce a compound of definite, reproducible composition. The decomposition is usually accompanied by the evolution of a volatile product such as H_2O , CO_2 , NH_3 , or SO_2 . Repeated ignitions followed by weighings will indicate the completion of the decomposition by giving a constant weight.

2. **Partial decomposition of desired product.** As the result of too high an ignition temperature or of the presence of unburned carbon from the charred filter paper, the product may undergo an undesirable partial volatilization or partial reduction. Partial

⁹ For complete tables of the density of air at various temperatures and pressures, see Treadwell and Hall: *Analytical Chemistry*, Eighth Edition, Vol. II, p. 14, Wiley, New York, 1938.

volatilization can usually be remedied by the addition of a reagent, the excess of which will volatilize, that will restore the volatilized portion. For example, PbSO_4 may be partially decomposed to PbO . A few drops of concentrated sulfuric acid may be added to the cooled residue, and, after another ignition at a more closely controlled temperature, the precipitate may be pure PbSO_4 with the excess of sulfuric acid volatilized. If partial reduction takes place, the product may be oxidized back to the desired form, either by strong heating with free access of air or by the addition of a volatile oxidizing agent such as nitric acid.

3. Volatilization of the crucible. Crucibles of porcelain or silica show no appreciable loss on strong heating, but crucibles of metal, such as platinum or platinum-iridium alloys, do show appreciable loss in weight on continued heating at high temperatures. Burgess and Waltenberg¹⁰ have shown that a crucible of 97.5 per cent Pt and 2.5 per cent Ir experienced the following loss in weight at the indicated temperatures:

$t^\circ \text{C}$	Loss (mg per 100 cm^2 per hour)
below 900	0
1000	0.57
1200	2.5

When a precipitate is ignited in a metal crucible, the weight of the precipitate may be corrected for crucible loss either by determining the rate of loss for the particular crucible under the conditions of the ignition and applying a correction factor or by weighing the empty crucible both before and after the ignition to obtain the correction for that particular experiment.

Weighings in gravimetric analyses are by no means confined to crucibles. Another important type of container that is often weighed is the absorption bulb or tube which is usually of glass. Of the likely sources of error in the weighing of glass objects, the following deserve consideration:

4. Variation in weight of adsorbed moisture. Containers made of glass, silica, or quartz are particularly susceptible to the formation of an adsorbed film of moisture on the surface. Glass is the worst offender of the three. The weight of the moisture varies widely with humidity and with the process of wiping or

¹⁰ Burgess and Waltenberg, *Ind. Eng. Chem.*, **8**, 487 (1916).

drying. Any correction of this phenomenon is, of necessity, of a highly empirical character. About the best generalization that can be made is that each laboratory should study its own particular conditions and adopt as a standard procedure a fixed method of wiping and a fixed time of standing outside and inside the balance case that has been found to yield reproducible results. It may be remarked that the custom of placing a dish of desiccant in a balance case has been shown to be about as effective as an incantation by a witch doctor from the jungles of Africa.

5. Electrification of glass surface in wiping. Rubbing a glass surface with cloth or chamois skin produces an electrostatic charge which, though it has no sensible weight, may result in an attractive or repulsive force on the balance pan amounting to as much as 100 mg in rare instances.¹¹ A standard time of, say, 10 or 15 minutes should be allowed to elapse between wiping and weighing to allow for dissipation of the charge. Various proposals of placing a specimen of radioactive material in the balance case to ionize the air and speed dissipation of the charge have not met with general favor.

The steps of an analytical procedure leading up to the final determination of the weight of a precipitate are laden with possibilities for error. Some of the errors are inherent to the analytical method employed; others are inherent to the manipulative processes involved. Some of the sources of error most frequently encountered in the course of analytical separations are:

6. Solubility of the precipitate. No substance is completely insoluble. The very nature of the solubility-product principle forbids total precipitation.¹² It is true that many substances are so slightly soluble that the amount dissolved in, say, 100 ml cannot be detected by ordinary weighing processes. Many of the compounds which are separated as gravimetric precipitates are soluble to an appreciable extent. If the solubility is known, the measurement of the volume of filtrate and wash solution may yield sufficient information for the computation of a correction term which must be added to the weight of the precipitate. In some instances, however, the time allowed for precipitation to take place does not permit the attainment of equilibrium, and this is especially true of the process of washing the precipitate. Unless equilibrium has

¹¹ Miller, *J. Am. Chem. Soc.*, **20**, 428 (1898).

¹² Yoe, *Chemical Principles*, p. 101, Wiley, New York, 1937.

been attained the amount dissolved may be widely variable. A more satisfactory method of correction consists in a direct determination of the dissolved constituent in the combined filtrate and washings by a nephelometric or colorimetric procedure. This method is used in atomic-weight determinations and analytical work of such a special nature. As a rule, the dissolving of a precipitate during washing can be minimized by using for the wash liquid a dilute solution which contains an ion in common with the precipitate. The repressing action in the common-ion effect can greatly retard the dissolving of the precipitate during washing. For some types of precipitates an organic wash liquid (e.g., alcohol, ether, acetone, etc.) is preferable.

7. Colloidal dispersion of the precipitated substance. Almost every procedure for regulating the conditions of precipitation must take account of the tendency of a great many substances to adsorb certain ions in the course of precipitate formation and become electrically charged molecular aggregates which are said to be colloidally dispersed. Even though the conditions of precipitation might have minimized colloid formation, the tendency may persist in the process of washing the precipitate. To overcome the possibility of error from this source, the wash solution must contain a small amount of an electrolyte. The electrolyte may be an acid, base, or salt, but whatever its chemical type, it must be volatile at the temperature of drying or ignition. The compounds most frequently used for this purpose are ammonium salts.

8. Contamination of precipitate. Only in isolated instances is the precipitated compound of a high degree of purity. The contamination may be caused in the following ways:

(a) *Simultaneous precipitation*, in which the precipitating ion is present in sufficient quantity to exceed the solubility product of an ion present as impurity, in addition to the solubility product of the ion it is desired to precipitate.

(b) *Formation of mixed crystals*, in which case the impurities are incorporated in the crystal lattice and do not change the regular structure of the latter.

(c) *Occlusion (real coprecipitation)*, in which case the impurities are not incorporated in the crystal lattice but are adsorbed during the growth of the crystals, causing imperfections in the crystals.

(d) *Surface adsorption* by the precipitate after it has been formed or separated. (Of practical importance only when the

precipitate has a large surface, i.e., when it behaves like a flocculated colloid.)

(e) *Post-precipitation*, which is a contamination of the precipitated compound after its formation, due to the slow formation of a foreign precipitate from its supersaturated solution.

(f) *Formation of a definite chemical compound*, e.g., the formation of a double salt. (Seldom encountered in analytical work.)

A full discussion of these various types of contamination and the means of eliminating the errors resulting from their respective actions is outside the scope of this book. For a clear understanding, the reader should consult the literature, especially the publications of Kolthoff¹³ and associates.

VOLUME

The calculation of the volume of a simple geometrical solid from linear measurements is one of the few examples of straightforward volume measurement in chemistry. Practically all other measurements of volume are actually comparisons of the capacities of vessels. It has already been pointed out in Chapter VI that the unit of volume (cubic centimeter) is different from the unit of capacity (milliliter). The errors of true volume measurement are the combined errors of the linear measurements upon which calculations are based. Problem 12 in Set XIV (p. 200) deals with the precision of volume calculation.

The volumes of gases, of liquids, and, in an indirect fashion, of irregular solids are obtained from capacity measurements. We shall deal only with the volumetric measurements of liquids in the succeeding remarks. The sources of corrigible error to which volumetric measurements of liquids are liable are: (1) use of uncalibrated vessels, (2) use of calibrated vessels presuming an erroneous standard temperature, (3) parallax in estimating the position of the meniscus, (4) non-reproducible liquid film in delivery vessels, and (5) use of volumetric vessels at temperatures other than the standard laboratory temperature. The means of correcting for these errors will be discussed in the order in which they are listed.

Calibration of Volumetric Ware. Capacity measurements are based upon the standard liter, which has already been defined

¹³ Kolthoff and Sandell, *op. cit.*, Chapters VII and VIII; Kolthoff, *J. Phys. Chem.*, **36**, 860 (1932); Kolthoff, *Z. anal. Chem.*, **104**, 321 (1936).

(p. 86) as the volume of pure water which at 4°C and 760 mm weighs a kilogram *in vacuo*. A vessel of glass which at 4°C has a capacity of a liter would have a different capacity at another temperature because of the expansion and contraction of glass with change of temperature. Inasmuch as 4°C is a highly impracticable temperature at which to carry out most volumetric measurements, a concession must be made to practical consideration. The National Bureau of Standards has adopted 20°C as the national *standard laboratory temperature* chiefly because it represents a reasonable, year-round average room temperature in this country. Consequently a glass volumetric vessel will have its designated capacity at 20°C and only at 20°C . The practical definition of a liter becomes: the volume occupied by that quantity of pure water at 4°C and 760 mm pressure which weighs a kilogram *in vacuo*, when the containing flask has a temperature of 20°C . The definition sounds complicated, but the complication is only apparent, not real, in laboratory practice. It will be necessary, before proceeding further, to differentiate between the two kinds of volumetric vessels. One kind is designed to **contain** a definite volume, and the other variety is designed to **deliver** a definite volume or equally definite fractions of that volume. The calibrations of the two varieties present slightly different problems; we shall discuss containing vessels first.

The weight of the water which, at temperature t , fills the flask to the mark is obtained by the difference in weights of the flask, clean and dry, and of the flask containing water. Water to be used in calibration work should be a good grade of distilled water which has been freshly boiled to remove dissolved gases (O_2 , N_2 , CO_2 , etc.) and cooled in a covered or loosely stoppered vessel. The weight must first be corrected for air buoyancy to give the *in vacuo* weight, making use of equation (10). To obtain the volume of the water from the weight, the weight must be divided by the density (expressed as grams per milliliter) for the temperature, t . The value of density should properly be in grams per milliliter¹⁴ or in terms of specific gravity,¹⁵ which is numerically the same. For most purposes the values of density in grams per

¹⁴ *J. Research, Nat. Bur. Standards*, **18**, 213 (1937), or *Smithsonian Physical Tables*, Eighth Edition, p. 166, 1934.

¹⁵ *Handbook of Chemistry and Physics* (cited on page 99), table of "Relative Densities of Water."

cubic centimeter are sufficiently exact to use in this computation. The volume obtained gives the capacity of the flask at the temperature t . This volume must be corrected for volume change if t is not 20°C . The coefficient of cubical expansion for glass varies from 0.000023 to 0.000028 per degree Centigrade. An average value which suits the purpose is 0.000025 per degree Centigrade. The capacity corrected to 20°C is given by

$$V_{20} = V[1 + 0.000025(20 - t)] \quad (11)$$

where V is the observed volume, V_{20} is the volume at 20°C , and t is the observed temperature. In ordinary quantitative analysis it is generally unnecessary to make a temperature correction unless the laboratory temperature differs by more than 2 or 3 degrees from the calibration temperature.

If it is desirable to have a flask contain correctly the nominal face value, the foregoing calculation can be reversed to yield a value for the weight of water, weighed in air, which must be placed in the flask. By noting the position of the meniscus, a new graduation mark may be engraved or denoted by a gummed strip.

The calibration of pipettes or burettes involves the delivery of water from the whole of, or between graduations of, the vessel into a suitable weighing bottle. The weight having been obtained, the volume of the water is calculated in exactly the same way as just described for a flask. In connection with delivery pipettes which deliver one definite volume, as distinguished from graduated pipettes, one further point must be noted in the calibration. The pipette is so constructed that, if the water is drawn up to the graduation mark, the pipette will then deliver the definite volume provided that the quantity of water remaining in the tip is reproducible. Among other things, the reproducibility is governed by whether the amount of water is the quantity held by true capillary equilibrium. To insure equilibrium, the tip of the pipette may be touched to the wall of the vessel into which delivery is being made. Then, with the water in the tip in contact with water outside the tip, equilibrium is readily attained. There are other methods, such as touching the surface of the liquid in the receiving vessel with the tip. The point is, whatever method of assuring equilibrium is employed in the calibration, that same method must be adhered to invariably in subsequent use of the pipette.

The volumetric instruments which can be obtained already calibrated from other countries, such as England or Germany, are very likely to have been calibrated for a standard temperature other than 20°C . Such vessels should either be recalibrated or a suitable allowance should be deduced to care for the discrepancy.

Parallax. In the discussion of calibration it has been presumed that the positions of menisci were correctly estimated. The extreme vagueness in the appearance of the top and bottom of the meniscus renders the recognizing of one or the other boundary very difficult. The meniscus may be brought into sharp relief by an arrangement (black and white card or sliding collar of black paper) to cast a dark shadow on the under-side of the concave meniscus. It is necessary that the shadow-casting object be at a fixed distance below the meniscus for all readings. With the meniscus in sharp relief, the next requirement is that the center of the eye be in the same horizontal plane with the meniscus; otherwise the meniscus will appear high or low on the engraved scale depending on whether the eye is below or above the horizontal plane. If the graduation marks are engraved completely or halfway around the glass cylinder, sighting points are provided. The proper sighting plane may be insured by means of a telescope sighting device such as is now manufactured.

Liquid Films in Delivery Vessels. For a pipette or burette to deliver a reproducible volume of liquid, the liquid film adhering to the inner wall must always be the same thickness. The thickness of the film depends upon several factors. These are: (1) *Cleanliness*. In all chemical work the vessels must be sufficiently clean to prevent contamination of the contents. For volumetric delivery vessels, supercleanliness is essential. This supercleanliness can be attained by the use of either a strong oxidizer like chromic acid or a mild alkali like trisodium phosphate, both of which can remove the last traces of "greasy" substances. Unless the glass surface is clean, the liquid cannot leave a uniform film. (2) *Time of outflow*. It is obvious that liquid flowing slowly will leave a thinner film than one flowing rapidly from a delivery vessel. This is true because of the viscosity of all liquids. For the formation of uniform films, an optimum time should be determined for each vessel which will result in a uniform rate, not of evacuation but of surface exposure. The National Bureau of Standards, as the result of an intensive study, has formulated requirements for the outflow times

of burettes and pipettes of various sizes.¹⁶ The time of outflow is determined by the diameter and shape of the outlet tip. The time of outflow of each delivery vessel should be measured to find whether it fulfills the Bureau's specifications. In the use of burettes and graduated pipettes, it is seldom that the outflow is unrestricted. To allow for the decreased rate and for intermittent flow, it is customary to allow one-half minute to elapse between cessation of flow and reading the position of the meniscus. This pause tends to allow the attainment of surface-film equilibrium.

(3) *Viscosity of liquid.* The film thickness depends upon the viscosity of the liquid being delivered; hence it follows that a pipette which has been calibrated with a liquid of a given viscosity can be used only with liquids having the same viscosity. Dilute aqueous solutions of electrolytes have sensibly the same viscosity as water. A pipette which was calibrated with water will not deliver its nominal volume of, say, concentrated sulfuric acid or ethyl alcohol. For the measurement, with a reasonably high degree of exactness, of volumes of liquids which have viscosities different from that of water, a pipette should be calibrated with the given liquid and reserved for use solely with that liquid.

Temperature Corrections. When volumetric ware has been calibrated for one temperature, its use for measuring liquids at other temperatures will result in error unless corrections are applied. The corrections must take account of the expansion or contraction of both liquid and vessel. As long as water is being measured, the correction is a simple matter, for we have available the values of the density of water over a wide range of temperature, and the glass vessel correction can be computed from equation (11). When the liquid being measured is an aqueous solution, not only is the density different from that of water but also the rate at which the density changes with temperature variation is different. Unfortunately, we do not have exhaustive tables of the densities at different concentrations and different temperatures for all the solutions likely to be encountered. The densities of many electrolyte solutions have been determined for wide concentration ranges, but unfortunately the ranges of temperature are very small.¹⁷

¹⁶ *National Bureau of Standards Circular 9*, 1916.

¹⁷ For one of the more extensive tables consult Landolt-Börnstein, *Physikalisch-Chemische Tabellen*, edited by Roth and Scheel, 5. Auflage, Table 87, Julius Springer, Berlin, 1923-1935.

Studies of the coefficients of cubical expansion of solutions of electrolytes have disclosed a surprising regularity in behavior for solutes of wide variety. It has been found that, in general, solutions of electrolytes of the same *normal* concentration have approximately the same coefficient of cubical expansion. The value being given for the coefficient of cubical expansion, h , for electrolyte solutions of definite strength, the value of the volume corrected to 20° C is given by

$$V_{20} = V_t + V_t h (20 - t) \quad (12)$$

in which V_t is the apparent volume at temperature t , and h is the value of the coefficient for the given normal concentration. The values of h for different concentrations expressed as milliliters per milliliter per degree Centigrade, are as follows:

CONCENTRATION	h
0 to 0.1 N	0.00020
0.5 N	0.00025
1.0 N	0.00029

These are average values which also make correction for the glass expansion. Very good approximate corrections ¹⁸ are obtained over the range 15° C to 25° C. For more exact corrections, pycnometric studies of the given solution need to be made. As mentioned on page 113, the temperature correction is usually not significant when the laboratory temperature is within 2 degrees of the calibration temperature, and it may be neglected except in work requiring the highest degree of precision.

CONGRUITY IN MEASUREMENTS

It was pointed out in Chapter VI that, in the measurement of a physical quantity, the proper procedure is to determine which of the variables that condition the system, of which the quantity is characteristic, are ponderable. By ponderable is meant whether the variable exerts an effect that is sufficiently appreciable to affect the value of the quantity within the desired number of significant figures sought. The ponderable variables must then be fixed as closely as possible, or allowance must be made for their effect by means of a correction factor.

¹⁸ In this connection see Kolthoff and Furman: *Volumetric Analysis*, Vol. II, p. 30, Wiley, New York, 1929.

Two other procedures, both of which are less desirable than the foregoing, are: (1) correct for everything in sight with no regard for the significance of figures; and (2) correct for nothing—let the variables fly at will, and let the result fall where it may. Procedure (1) is that of an experimenter of great skill, with a brain full of facts but imperfect logical capacities; procedure (2) is that of a bungler. The proper course is somewhere between these two extremes.

The extent to which the correction of an apparent measured value is carried will depend upon what use is to be made of the value. Examples will illustrate. In the calculation of the volume of a gas at S.T.P. from the observed volume at temperature t and pressure p by means of the ideal gas law, p must be measured barometrically. If the volume were known to four significant figures, the apparent value of p should be corrected by the application of such of the known corrections as affect the fourth figure of the value of p . Those corrections which have no effect on the fourth figure should be neglected as imponderable. On the other hand, suppose that the observed value of p is to be used in the computation of the correction for the boiling point of water by equation (1), which is,

$$t = 100 + 0.0375 (b - 760)$$

If it were desired to know the boiling point to only the nearest tenth of a degree, none of the known corrections would need to be applied. The best rule to follow in applying corrections in measurement is simple common sense as embodied in the rules for significant figures.

In most measurements, the final result is computed from the results of several separate observations which constitute the procedure of measurement. We must reserve for Chapter X the deduction of a basis for the relationship between the error of a computed result and the separate errors of the values entering the computation, but as a general rule the various steps of the procedure should be executed with an approximately equal degree of exactness and carefulness. A striking example of incongruity in measurement was observed recently by one of the authors. (J.H.Y.) A student in determining the normality of a base by titration against a carefully standardized acid solution was found to be titrating a measured volume of the acid with the base solution delivered from

a calibrated burette to which temperature corrections were applied, *but he had measured the solution of acid with an ordinary 50-ml graduated cylinder.* It is hardly necessary to add that this student was a beginner in quantitative analysis.

In the practice of the art of laboratory research, more time is wasted on unjustified refinements and unnecessary corrections than on any other one item that could be named. The inhabitants of certain rural areas of the Southern states often voice a bit of folk-wisdom which exemplifies this kind of incongruity. They speak of an unaccustomed accessory or elegant frill as being "like a fifty-dollar saddle on a twenty-five-dollar horse."

It should not be understood from the foregoing remarks that we decry the careful experimenter who is sufficiently devoted to the cause of science to indulge a few extra pains in order to satisfy his own feeling of work well done. If no one had a desire to improve existing theories and refine old measurements by taking a few extra pains, there could be no progress in science. The cause for which we plead is not a stagnation through mediocrity but a rapid progress through consistency.

RANDOM ERRORS

To the class of random errors belong primarily those errors "about which we can do nothing." They are called random because they are attributed to the free and practically independent operation of those variables of a system which have been deemed individually inappreciable. Free and independent operation would result in errors which in magnitude would be scattered at random.

Most of the corrections which we have considered at length are approximate measurements or approximate fixings of variables on the borderline between ponderable and imponderable. Hence, the application of imperfect corrections results in **residual errors** in the measurements to which the corrections are applied. Although the refinement of corrections tends to minimize these residual errors, even the refined corrections will by their very nature be subject to random errors. We recognize a kinship between random and residual errors and therefore class them together provisionally under the title of "random."

In practically every measurement we use instruments which have been carefully constructed, but there is a limit to the fineness

with which a scale can be graduated. It would appear that the custom of estimating distances between the finest graduations, to add just one more figure to the observed value, must have grown either out of a sporting instinct or perhaps from a miserly sort of grasping instinct. At any rate, the estimation between scale divisions is the final step and an integral part of all instrumental measuring procedures. Naturally, estimation results in uncertainty, hence there is an error in the estimated value of the last figure. Manifestly the **errors of estimation** are inherent in the instrumental method, and no control over them can be exercised other than minimization through practice in estimation. With some few types of scales, the vernier principle may be applied, but even then some of the more zealous attempt the estimation of distances between the graduations of the vernier and of the scale. Although it is possible through repeated observations to determine the numerical value of a precision measure for uncertainty of estimation, the application of a correction is out of the question, and the errors of estimation are of necessity classed with random errors.

Personal Errors. The goal of any branch of science as it becomes more "scientific" is the attainment of highly objective methods of experimentation in preference to subjective methods. An ideally objective method would be one in which the judgments or prejudices of the experimenter could have no part. A completely subjective method would depend entirely upon the mental processes, either conscious or unconscious, of the observer. Even the attempt to define objectivity introduces subjectivity; hence we could hardly expect to attain a completely objective mode of experimentation. On the other hand, the postulation of a completely subjective method implies that it is objectively subjective. An experimental procedure may be more objective than subjective or the reverse. It can never be completely one or the other. The progress of a science is marked by the elimination of subjective aspects of experimentation and a consequent increase in the degree to which the methods are objective. In recognizing that the subjective element can never be removed, we are forced to the conclusion that in any scientific measurement there is at least one step or observation which depends upon sensory perception.

Many writers on the subject of errors speak somewhat hastily of the **personal equation** involved in scientific measurement. We are going to hazard a guess about the nature of the personal equa-

tion. For information on the subject of perception we must turn to the psychologists. Weber (1796–1878) in 1834 formulated his psychophysical law of sensory discrimination. Weber's law can be stated symbolically:

$$\frac{\Delta I}{I} = \pm K \quad (13)$$

The meanings of the symbols are as follows: I represents intensity of stimulus; ΔI represents the just noticeable difference in intensity of stimulus from the preceding stimulus; and K is a constant. A concrete example will probably serve better to explain than involved definitions. It has been found that, on the average, a person with normal eyesight can, in 75 per cent of the trials made, recognize a difference in intensity between two lights that is related to the intensity of the first light by the ratio 1/100. If the original stimulus is an illuminated field with a brightness of 10.0 millilamberts, the brightness that can be recognized as greater in 75 trials out of 100 is 10.1 millilamberts. Or, for an original brightness of 20.0 millilamberts, the brightness just noticeably less would be 19.8. In the first case,

$$\frac{10.1 - 10.0}{10.0} = \frac{1}{100} = K_1$$

and in the second case,

$$\frac{19.8 - 20.0}{20.0} = -\frac{1}{100} = K_2$$

Weber's law has been shown to be generally true in several sense departments,¹⁹ but the value of K is not the same for two different sense departments. The following tabulated values ²⁰ of K are for several different senses.

SENSE DEPARTMENT	$\pm K$
Visual.....	0.01
Muscular (weight lifting).....	0.025
Skin pressure.....	0.05
Auditory.....	0.20

¹⁹ For an application of Weber's law to colorimetric analysis, see Yoe and Wirsing, *J. Am. Chem. Soc.*, **54**, 1866 (1932); see also Horn and Blake, *ibid.*, **36**, 516 (1906).

²⁰ The data in the table are computed from values given in Dashiell: *General Psychology*, p. 254, Houghton Mifflin, Boston, 1937.

It should be emphasized that the above values of the ratio are the results for trials performed under optimum conditions for sensory discrimination. For conditions other than optimum, the value of the ratio for a given sense department will be different from the commonly assigned value, but it will be constant for a particular set of conditions.

Perception is subject to a further set of limitations. Below a threshold value of intensity of stimulus, sensation is not evoked when a stimulus is presented. In addition, above an upper limit sensory discrimination becomes impossible. Thus a light may be too dim to be seen or so bright that it produces glare. The experimental psychologists have demonstrated that the threshold value remains remarkably constant for different observers, provided as before that the conditions of the test are optimum. It has been found that Weber's law holds best at about the middle of the range of stimulus intensities between the threshold and the upper limit. On each side of this optimum, middle range the ratio invariably becomes greater. The commonly assigned values of the ratios represent minima or limiting values. The accompanying graph, Fig. 8, shows the variation of the ratio for light intensities. It will be observed that within the range of 1 to about 500 millilamberts, which is around the middle of the total range, the ratio remains fairly constant.

Most psychologists agree that Weber's law is not universally valid. To avoid embroilment in a debate, we shall only postulate that Weber's law represents a general limiting tendency in perception. The necessary conclusion to be drawn from our postulate is that, since only finite increments of stimuli can be recognized, then the personal error involved in the subjective steps of a measurement will have a finite value. Further, the personal error, excepting the small mistakes due to bias, will tend toward a constant limiting value for a particular individual in the given method of measurement. Allowance must be made for a slight improvement which may result from an observer's conscious or unconscious approach toward optimum conditions.

Granted the foregoing argument, we have established a possible basis for the frequently occurring statement that "personal errors are usually constant." All stimuli having "true" values between $I - \Delta I$ and $I + \Delta I$ would be recognized as I . If a sufficient number of trials were made, we should expect the average

error to approach $\frac{\Delta I}{2}$, although the assurance would not be very great. We should hesitate to recommend that any experimenter attempt to evaluate his personal error from Weber's law and apply corrections derived therefrom. A complete representation of the personal equation would have to account for many items in addition to sensory discrimination. It would have to deal with reac-

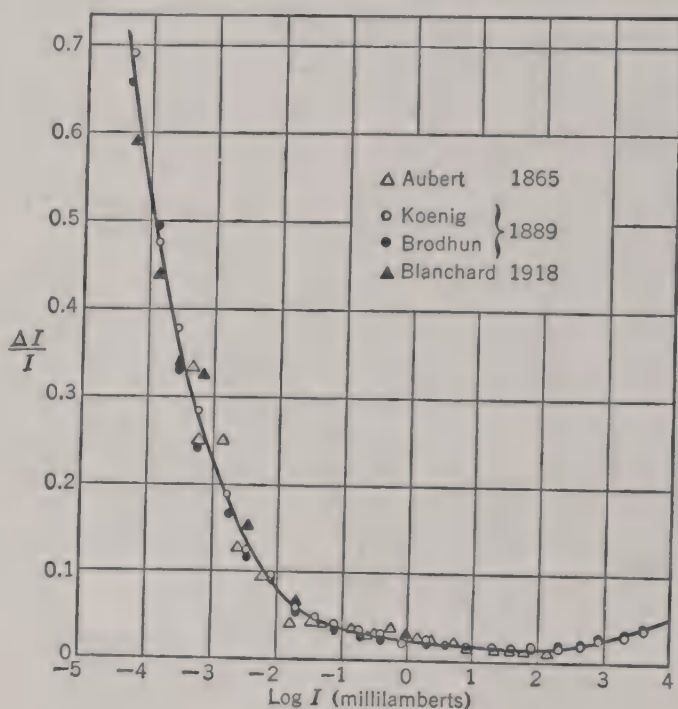


FIG. 8.

By permission from "A Handbook of General Experimental Psychology," C. Murchison, Editor, Clark University Press, Worcester, Mass., 1934. From Chapter 14 by Selig Hecht.

tion time, tendency to preconceived judgments, fatigue, etc. The general tendency exemplified in Weber's law is a necessary part but not the complete representation of the personal equation.

In the sense that personal errors may be slightly lessened by accommodation, they are partially corrigible. In the sense that personal errors are finite and ever-present, they belong to the class called random errors.

PROBLEM SET XI

1. When the barometric pressure is 712.5 mm, what is the boiling point of water, computed to the nearest hundredth of a degree?

2. At this same pressure of 712.5 mm, what is the boiling point of ethyl alcohol?

3. Two Centigrade thermometers, one uncalibrated and the other a thermometer certified by the National Bureau of Standards, are each submerged to the 15° mark in a thermostat. The following data were obtained:

	UNCALIBRATED	CERTIFIED
Reading	72.07°	71.94°
Mean stem temperature	29°	26°

For the temperature 70° C, the certificate gives a correction of +0.02°. Compute (a) the temperature of the thermostat and (b) the correction term for the uncalibrated thermometer at that temperature.

4. At 27.5° C, the reading on a brass-scale barometer was 752.38 mm. The diameter of the barometer tube is 8 mm, and the height of the meniscus is 0.6 mm. The correction for capillary depression is +0.35 mm. Compute the atmospheric pressure.

5. A certain brass object was placed on the right pan of a balance, and the sum of the weights required to counterbalance it was 26.4824 g. The same object was then placed on the left pan and the apparent weight was 26.4882 g. Compute (a) the balance-arm ratio and (b) the "absolute" weight of the object, assuming the weights to have been "absolutely" calibrated.

6. The Gauss method of double weighing consists in observing the apparent right- and left-pan weights and taking the average of the two as the corrected weight. From the data of Problem 5, compute the average weight and compare with the answer obtained in 5(b).

7. A flask made of glass ($d = 2.4$) weighs 12.3821 g in air when counterbalanced by brass weights ($d = 8.5$). Determine the weight of the flask *in vacuo*.

8. The *in vacuo* weight of a certain platinum crucible ($d = 21.4$) is 19.5293 g. What would this crucible weigh in air against weights made of quartz ($d = 2.66$)?

9. These data were obtained in the calibration of a 250-ml volumetric flask:

Weight of flask (empty)	56.882 g
Weight of flask (filled to mark with water).	305.276 g
Temperature of water	25.4° C
Density of water at 25.4° C	0.9969718 g per ml

Weighings were made with brass weights ($d = 8.5$) on a balance for which the value of the ratio, R/L , is 1.00002 at 50-g load and 1.00019 at 300-g load. Compute the capacity of the flask. Standard laboratory temperature is 20° C.

10. In the standardization of an approximately 0.05 *N* solution of carbonate-free sodium hydroxide against a 0.04872 *N* solution of hydrochloric acid

with phenolphthalein as indicator, the following data were obtained using calibrated burettes:

	NaOH	HCl
Final burette reading.....	43.04 ml	44.86 ml
Initial burette reading.....	0.16 "	0.07 "
Burette correction.....	+0.03 "	-0.08 "
Temperature of solution.....	26.1° C	25.8° C

Compute the normality of the sodium hydroxide solution at 20° C on the basis of this single titration.

CHAPTER VIII

STATISTICAL METHODS

Through and through the world is infected with quantity. To talk sense, is to talk in quantities. It is no use saying that the nation is large,—How large? It is no use saying radium is scarce,—How scarce? You can not evade quantity. You may fly to poetry and to music, and quantity and number will face you in your rhythms and octaves. Elegant intellects which despise the theory of quantity are but half developed. They are more to be pitied than blamed.

—A. N. WHITEHEAD.

Statistics is concerned with counting, classifying, and interpreting results. We shall be concerned with statistics only in the third aspect, namely, the interpretation of series of numerical values which are discordant among themselves. Statistics has been defined as the science of averages; this definition will suit our purpose.

AVERAGES

An average is a measure of central tendency. The necessary criterion of an average, if it is to represent central tendency truly, is that it must have a value such that, if all the measurements in the group became equal, each would equal the average. Several kinds of averages are in use, each of which is best for a certain type of statistical group.

Mode. The mode is that measurement of a group which occurs most often. It is the most likely or most probable value in the particular group. An example of the use of the mode in sociological statistics is in evaluating the average income of the residents of a geographical unit, say a county. The accidental presence of a millionaire on a country estate in that county would give a disproportionate estimate to the per capita wealth of the residents in comparison with the adjoining county which had no millionaires. The mode would disregard the millionaire and take

into account only the regular residents who purchase vacuum sweepers, sewing-machines, baby carriages, etc. In other words, the mode as an expression of central tendency is less susceptible to accidentally extreme variations.

Median. The median is that measurement in a group above which exactly half of the values fall and likewise below which exactly half fall. If a series of measurements is arranged in order of increasing magnitude, the median is the middle measurement. The median is of value in that it is simple to obtain and it is almost completely independent of extreme variations. Its use in statistics is limited because it is dependent only upon the number of measurements, and unless the values happen to be distributed in the most usual manner, the median will give a poor representation of central tendency.

Arithmetic Mean. The most commonly used average is the arithmetic mean. Certainly, in chemical measurements, it is agreed that the arithmetic mean gives the most nearly correct representation of central tendency. The arithmetic mean was the average meant in the discussion of "best value" in Chapter VI. The definition of the arithmetic mean is the sum of all the measurements divided by the number of measurements. In the language of mathematics,

$$a = \frac{x_1 + x_2 + x_3 + x_4 + \cdots + x_n}{n} \quad (1)$$

where the x 's represent the individual measurements. Or in simpler form,

$$a = \frac{\Sigma x}{n} \quad (2)$$

where Σ here, and as hereinafter used, denotes the "sum of." An important relation following from (2) will be used later. It is:

$$\Sigma x - na = 0 \quad (3)$$

Root Mean Square. Another type of average that is worth brief notice is the root mean square average. By definition it is:

$$\text{rms} = \sqrt{\frac{\Sigma x^2}{n}} \quad (4)$$

It might be called the square root of the arithmetic mean of the

squares of the values in the group. A familiar example of the use of this average is found in the kinetic theory of gases. In energy calculations, it is the average of the squares of the velocities that is proportional to the average energy, since the energy is proportional to the square of the velocity of a particle (atom or molecule). On the other hand, in calculations of the collisions between particles, the average velocity used is the arithmetic mean.

Weighted Arithmetic Mean. When a series of measurements has been compiled from experiments performed under different conditions (i.e., made by different observers or by one observer under conditions not uniformly favorable) it is often the practice to give greater weight to the more favorably determined values. Under these circumstances a modified form of the arithmetic mean is used. The modified form is known as the weighted arithmetic mean which is defined as:

$$\text{w.a.} = \frac{w_1x_1 + w_2x_2 + w_3x_3 + \cdots w_nx_n}{w_1 + w_2 + w_3 + \cdots w_n} \quad (5)$$

Or, more simply,

$$\text{w.a.} = \frac{\sum wx}{\sum w} \quad (6)$$

The assignment of weights is of necessity a highly arbitrary procedure, and the practice of weighting must be performed with great caution. It requires an experimenter of wide experience and with a fine sense of distinctions to assign weights in such a way that any widespread confidence in the result will be inspired. In Chapter X we shall derive a basis for weighting the averages of measurements by different observers. For a set of observations made by one observer under varying conditions of favorability, no criterion of weighting has been evolved. The most plausible course is to take more observations.

DEVIATIONS

Since the arithmetic mean is the average that will be used almost exclusively, deviations from a will be the deviations in which we shall be most interested. By convention, the deviation is always found by subtracting the mean from the observation. Expressed in symbols,

$$d_a = x - a \quad (7)$$

An important property of the deviations from a is that the sum of the deviations from a in a group is zero. Assume a group of n measurements of which a is the arithmetic mean.

$$\begin{array}{rcl}
 d_{a_1} & = & x_1 - a \\
 d_{a_2} & = & x_2 - a \\
 d_{a_3} & = & x_3 - a \\
 . & & . \\
 . & & . \\
 . & & . \\
 d_{a_n} & = & x_n - a \\
 \hline
 \Sigma d_a & = & \Sigma x - na
 \end{array} \tag{8}$$

By equation (3), the right-hand side of the expression is zero; hence

$$\Sigma d_a = 0 \tag{9}$$

Short-Cut Method for Obtaining the Arithmetic Mean. When the arithmetic mean of an extended series of measurements having a large number of significant figures must be obtained, the method of adding the large column of figures and dividing by the number of measurements is likely to prove tedious. Tedium produces fatigue and hence increases the possibility of mistakes in arithmetic. A method involving simpler computations and the handling of fewer digits is based upon the properties of the deviations from an arbitrary value x' . Assume n measurements in the group, and take an arbitrary value x' which lies near the middle of the series. Defining the deviation as:

$$d_{x'_1} = x_1 - x' \tag{10}$$

then, for the whole group,

$$\begin{array}{rcl}
 d_{x'_1} & = & x_1 - x' \\
 d_{x'_2} & = & x_2 - x' \\
 d_{x'_3} & = & x_3 - x' \\
 . & & . \\
 . & & . \\
 . & & . \\
 d_{x'_n} & = & x_n - x' \\
 \hline
 \Sigma d_{x'} & = & \Sigma x - nx'
 \end{array} \tag{11}$$

Dividing both sides of the summation equation by n ,

$$\frac{\Sigma x}{n} - \frac{nx'}{n} = \frac{\Sigma d_{x'}}{n} \quad (12)$$

Substituting from equation (2) and rearranging,

$$a = x' + \frac{\Sigma d_{x'}}{n} \quad (13)$$

Example 1. Successive determinations of the normality of a solution of ceric ammonium sulfate titrated against pure ferrous ammonium sulfate, using *o*-phenanthroline ferrous complex as an internal indicator, gave the following values: 0.1327; 0.1325; 0.1329; 0.1326; 0.1328; 0.1327; 0.1326; 0.1325; 0.1328; 0.1326; 0.1327.

Arrange the values in tabular fashion, and take as an arbitrary central value 0.1326, i.e., $x' = 0.1326$. In a second column set the deviations from x' and take the algebraic sum of the deviations.

x	$d_{x'}$
0.1327	+0.0001
.1325	- 1
.1329	+ 3
.1326	0
.1328	+ 2
.1327	+ 1
.1326	0
.1325	- 1
.1328	+ 2
.1326	0
.1327	+ 1
	+0.0008 = $\Sigma d_{0.1326}$

$$a = 0.1326 + \frac{0.0008}{11}$$

$$a = 0.1326 + 0.000073$$

$$a = 0.13267$$

According to Rule 9 (Chapter II), five figures may be carried in the average.

The arbitrary value need not be a value from the group. It may be a conveniently rounded number falling either between two values of the

group or outside the group. If, instead of 0.1326, the value 0.1330 had been taken, the same result would have been obtained. Thus,

$$\begin{aligned}\Sigma d_{0.1330} &= -0.0036 \\ a &= 0.1330 + \frac{(-0.0036)}{11} \\ a &= 0.13267\end{aligned}$$

For convenience, to make the algebraic sum of the deviations as small as possible, it is customary to take for x' a value near the middle of the group of measurements.

DISPERSION

When the values in a series of measurements of a quantity are scattered over a wide range, the measurements are said to be highly dispersed. Another way of stating the same thing is to say the group of measurements has high dispersion; or, conversely, a set which is closely grouped has low dispersion. For purposes of comparison we need a more rigorous statement than is afforded by the adjectives "high," "medium," or "low." Several modes of expressing dispersion have been developed by the statisticians.

Range. The simplest but least exact expression of dispersion is a statement of the range over which the values of the group vary. Thus, in Example 1, the highest value is 0.1329 and the lowest value is 0.1325. The range of variation is:

$$\text{Range} = 0.1329 - 0.1325 = 0.0004$$

The value 0.0004 by itself tells little unless it is coupled in some way with the magnitude of the quantity of which it is an attribute. Hence, a coefficient of dispersion must be defined. The coefficient of dispersion is given as

$$\text{Coefficient of dispersion} = \frac{\text{Dispersion measure}}{\text{Arithmetic mean}}$$

The range coefficient is therefore

$$\kappa_r = \frac{\text{Range}}{a} \quad (14)$$

Continuing with the data of Example 1,

$$\kappa_r = \frac{0.0004}{0.13267} = 3 \times 10^{-3}$$

The great weakness of the "range" as a deviation measure is that it is dependent solely upon the two extreme values and is independent of the number of measurements. Suppose that only two measurements of the normality of the ceric ammonium sulfate solution had been made and the values were 0.1329 and 0.1325. The arithmetic mean is 0.1327, and the dispersion coefficient would be

$$\kappa_r = \frac{0.0004}{0.1327} = 3 \times 10^{-3}$$

The final answer is the same in the two instances although the cases are widely different. We require therefore a more exacting measure of dispersion.

Average Deviation. A better measure of the dispersion of a group of values is given by the average of the deviations from the arithmetic mean. This gives a dispersion measure which is a function of all the values and not just of the extremes. The average deviation (δ) is defined as

$$\delta = \frac{\sum |d_a|}{n} \quad (15)$$

where the vertical bars enclosing the symbol for deviation mean "the absolute value of," disregarding plus and minus signs. As before, the coefficient of dispersion is defined as

$$\kappa_\delta = \frac{\delta}{a} \quad (16)$$

The average deviation has a weakness in that it tells nothing about the manner in which the values are dispersed. It is conceivable that a set of deviations all of medium size might accidentally have the same average as another set in which there were several very small, a few medium, and several very large deviations. The average deviation does not give any indication of the distribution of the deviations. In other words, the arithmetic mean of the deviations provides neither penalty for large deviations nor reward for small deviations.

Standard Deviation. The root mean square average of the deviations from the arithmetic mean, known as the standard deviation, is not only a function of all the values of a group but it also is more sensitive than the average deviation to the presence of

large or small deviations. Hence the standard deviation is the better measure of consistency. The standard deviation is defined as

$$\sigma = \sqrt{\frac{\sum d_a^2}{n}} \quad (17)$$

Squaring a number with either positive or negative sign gives positive values; hence it is not necessary in this formula to denote the absolute value of the deviations. Because the standard deviation is not subject to the faults of the other dispersion measures it is the most generally used. Statisticians agree that it gives the best representation of the dispersion of a group of measurements. The coefficient as in the preceding cases is given by

$$\kappa_\sigma = \frac{\sigma}{a} \quad (18)$$

Short-Cut Method for Obtaining the Standard Deviation. In a manner entirely analogous to that employed for the short-cut method of computing the arithmetic mean, the standard deviation may be obtained from the deviations from an arbitrary value (x'). First, we shall define the root mean square deviation from an arbitrary value as

$$s = \sqrt{\frac{\sum d_{x'}^2}{n}} \quad (19)$$

Let

$$a - x' = \bar{x}' \quad (20)$$

From equation (13) it follows that

$$\bar{x}' = \frac{\sum d_{x'}}{n} \quad (21)$$

By definition

$$d_{x'_1} = x_1 - x'$$

Therefore

$$s^2 = \frac{\sum (x - x')^2}{n}$$

From (20)

$$x' = a - \bar{x}'$$

Now

$$s^2 = \frac{1}{n} \sum (x - a + \bar{x}')^2$$

Carrying out the squaring in a separate operation

$$\begin{aligned}(x - a + \bar{x}')^2 &= [(x - a) + \bar{x}']^2 \\ &= (x - a)^2 + 2\bar{x}'(x - a) + \bar{x}'^2\end{aligned}$$

Replacing in the equation

$$s^2 = \frac{1}{n} \Sigma[(x - a)^2 + 2\bar{x}'(x - a) + \bar{x}'^2]$$

The sum of the terms enclosed by the bracket would be:

$$\begin{array}{cccc} (x_1 - a)^2 & + & 2\bar{x}'(x_1 - a) & + \bar{x}'^2 \\ (x_2 - a)^2 & + & 2\bar{x}'(x_2 - a) & + \bar{x}'^2 \\ \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot \\ (x_n - a)^2 & + & 2\bar{x}'(x_n - a) & + \bar{x}'^2 \\ \hline \Sigma(x - a)^2 & + & 2\bar{x}'\Sigma(x - a) & + n\bar{x}'^2 \end{array}$$

This becomes

$$s^2 = \frac{\Sigma(x - a)^2}{n} + \frac{2\bar{x}'\Sigma(x - a)}{n} + \bar{x}'^2$$

Then, from equations (17) and (9),

$$s^2 = \sigma^2 + \bar{x}'^2$$

or

$$\sigma = \pm \sqrt{s^2 - \bar{x}'^2} \quad (22)$$

In terms of deviations from the arbitrary value (x'),

$$\sigma = \pm \sqrt{\frac{\Sigma d_{x'}^2}{n} - \left(\frac{\Sigma d_{x'}}{n}\right)^2} \quad (23)$$

By means of a combination of the two short-cut methods that have been derived, it is now possible to compute both a and σ in one grand operation.

Example 2. In the analysis of a certain organic compound by the usual method of combustion and absorption of the water formed in a tube charged with a suitable desiccant, the following percentages of hydrogen were found: 2.71, 2.76, 2.79, 2.78, 2.82, 2.76, 2.78, 2.74, 2.76, 2.74.

For x' take 2.75; $n = 10$

x	$d_{x'}$	$d_{x'}^2$
2.71	-0.04	0.0016
2.76	+ 1	1
2.79	+ 4	16
2.78	+ 3	9
2.82	+ 7	49
2.76	+ 1	1
2.78	+ 3	9
2.74	- 1	1
2.76	+ 1	1
2.74	- 1	1
	+0.14	0.0104

By equation (13)

$$a = 2.75 + \frac{0.14}{10} = 2.764$$

And by equation (23)

$$\sigma = \pm \sqrt{\frac{0.0104}{10} - \left(\frac{0.14}{10}\right)^2}$$

$$\sigma = \pm \sqrt{0.00084}$$

$$\sigma = \pm 0.029$$

DISTRIBUTION

The measures of dispersion that have just been discussed are capable of yielding a numerical quantity which is characteristic of the degree of scattering of the values in a group. These measures, however, tell little of the manner of distribution of the values. Distribution is differentiated from dispersion in that distribution takes account of the order of occurrence and the frequency of occurrence in a group as well as the extent of scattering. We shall need to know something of the distribution of deviations before we can evolve numerical measures of precision which, in a limited sense, are approximate measures of accuracy.

Let us imagine a gigantic creature from another planet who has come to visit the earth. To him man would appear a biological type with certain attributes—height, weight, age, wealth, etc.—all of which are expressible as numbers. He would undoubtedly recognize that these attributes are susceptible of variation, but he

would still consider man as a type and not as a large number of highly diverse individuals. Take, for instance, the attribute of chest measurement. If our colossus observed the chest measurements of a large number of specimens he would note a variation. If he possessed intelligence in proportion to his size, he would ascribe the variation to a large number of variables: the amount of exercise taken by the specimen, the kind of food eaten, heredity, etc. He would further recognize that heredity is dependent upon a very large number of remote ancestors each of whom contributed some small share to the deviation of this specimen from an average value. To the giant the measurement of chests would be a "scientific measurement" comparable to the experimental work of the chemist.

The following table of the chest measures of 5,732 Scotch soldiers¹ will serve as an example.

CHEST MEASURE IN INCHES	NUMBER OF SOLDIERS
33	3
34	19
35	81
36	189
37	409
38	753
39	1,062
40	1,082
41	935
42	646
43	313
44	168
45	50
46	18
47	3
48	1

The first thing to be noted is that, in going down the table, there is an increase in the number of men having each "chest measure"; at about the middle a maximum occurs, and from there on a decrease is seen. For a given value, say 36 inches, the number of men having that value, 189, gives the frequency of occurrence of that value. It is apparent that the frequencies show a definite

¹ *Edinburgh Medical Journal*, 13, 260 (1817); cf., Whittaker and Robinson: *Calculus of Observations*, p. 165, Blackie, London, 1924.

and orderly trend. A graph will show the trend more clearly. Fig. 9 depicts an extremely usual state of affairs. This type of distribution, known as the normal frequency distribution, is the kind most frequently found in experimental science.

One further point concerning the table should be noted. At first glance it would appear either that (1) the measurements had been made only to the nearest inch, or (2) the men observed all

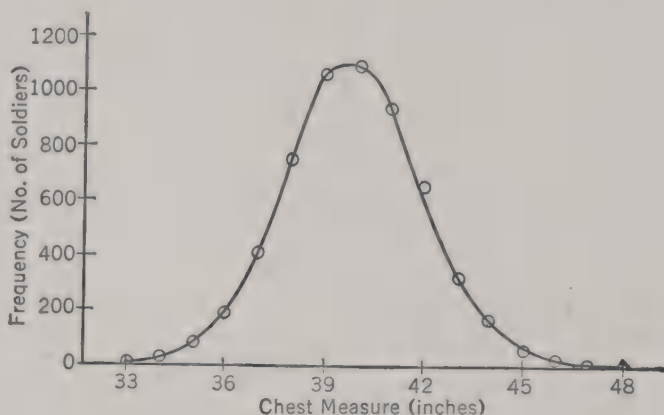


FIG. 9.

had chest measurements that were exact multiples of an inch, or (3) all the men having chest measures within a certain range had been grouped together. Undoubtedly the third possibility is the true one. The men have been divided into subgroups or classes, each class covering an interval of one inch, and the value stated is the mid-value of the class. Thus, in the first class, 3 men fell within the range $32\frac{1}{2}$ to $33\frac{1}{2}$; the next class is between $33\frac{1}{2}$ and $34\frac{1}{2}$, and so on down the table. This conception of **class interval** is a necessity because there is only a finite number of observations. The class interval might be made smaller, but since the measurements were made with only a limited degree of exactness, irregularities would appear that would obscure the order and produce chaos. It is only when an infinite number of measurements have been made that an infinitesimal interval (dx) can be taken. The class interval is a concession to practical considerations.

With the introduction of frequencies, the computation of averages and dispersion measures is a little more complex than in the previous instances. In the long run, however, for large numbers of values which can be grouped into classes, the labor is much less,

as can be seen by imagining a table containing 5,732 terms. Using the symbol f for frequency, the computation of a and σ for the chest measure distribution of Scotch soldiers is as shown.

Example 3. Taking 40 as x' ,

1	2	3	4	5
x	f	$d_{x'}$	$fd_{x'}$	$fd_{x'}^2$
33	3	-7	- 21	147
34	19	-6	- 114	684
35	81	-5	- 405	2,025
36	189	-4	- 756	3,024
37	409	-3	-1,227	3,681
38	753	-2	-1,506	3,012
39	1,062	-1	-1,062	1,062
40	1,082	0	0	0
41	935	1	935	935
42	646	2	1,292	2,584
43	313	3	939	2,817
44	168	4	672	2,688
45	50	5	250	1,250
46	18	6	108	648
47	3	7	21	147
48	1	8	8	64
5,732			-866	24,768

Note:

$$\Sigma f = n$$

$$\Sigma fd_{x'} = \Sigma d_{x'}$$

$$\Sigma fd_{x'}^2 = \Sigma d_{x'}^2$$

Note also:

Column 4 is the product of 2 and 3.

Column 5 is the product of 3 and 4.

$$a = x' + \frac{\Sigma d_{x'}}{n}$$

$$a = 40 + \frac{(-866)}{5,732} = 40 - 0.15$$

$$a = 39.8$$

And

$$\sigma = \sqrt{\frac{\Sigma d_{x'}^2}{n} - \left(\frac{\Sigma d_{x'}}{n}\right)^2}$$

$$\sigma = \sqrt{\frac{24,768}{5,732} - (-0.15)^2}$$

$$\sigma = \pm 2.1$$

Therefore

$$\kappa_{\sigma} = \frac{2.1}{39.8} = 0.053$$

PROBLEM SET XII

Note carefully: A complete record of the computations and answers to problems must be kept in order to solve succeeding illustrative examples in this and subsequent chapters.

1. The following determinations of the atomic weight of carbon by means of combustion analysis are from Baxter and Hale:²

12.0080	12.0101	12.0106
.0090	.0101	.0106
.0090	.0102	.0107
.0095	.0102	.0111
.0095	.0102	.0113
.0096	.0105	.0116
.0097	.0106	.0118
.0101	.0106	.0120
.0101	.0106	.0129

Determine the (a) mode, (b) median, (c) arithmetic mean, and (d) standard deviation.

2. In the determination of standard electrode potentials, Bates and Vosburg³ obtained the following values of E_{25}° for the Hg, Hg₂I₂, I⁻ electrode:

-0.04047	-0.04054	-0.04036
.04044	.04037	.04041
.04042	.04057	.04023

For these data obtain (a) κ_r , (b) κ_{δ} , and (c) κ_{σ} .

3. Linhart⁴ found these values of E_{25}° for the Hg, Hg₂⁺⁺ electrode:

-0.7896	-0.7929
.7895	.7922
.7927	

Compute (a) κ_r , (b) κ_{δ} , and (c) κ_{σ} . Which set of observations is more consistent: Linhart's or Bates and Vosburg's (Problem 2)?

4. Through a determination of the weight of silver chloride yielded by benzoyl chloride, Scott and Hurley⁵ found these values for the atomic weight of carbon:

12.0100	12.0105
.0102	.0099
.0103	.0101
.0101	.0106

² *J. Am. Chem. Soc.*, **59**, 506 (1937).

³ *Ibid.*, **59**, 1188 (1937).

⁴ *Ibid.*, **38**, 2356 (1916).

⁵ *Ibid.*, **59**, 1905 (1937).

Compute a and σ . Compare these data with those of Baxter and Hale (Problem 1) for consistency.

5. Lingane and Larson ⁶ measured the temperature coefficients of cells of the type:

Pt | Quinhydrone (s), $\text{HClO}_4(C_1)$ | $\text{HClO}_4(C_1)$, $\text{AgNO}_3(C_2)$ | Ag
at various acid concentrations. From their measurements they obtained the following values of the temperature coefficient of the standard silver electrode at the indicated acid concentrations:

C_{HClO_4}	$\left(\frac{dE_{\text{Ag}}^{\circ}}{dT}\right)_{298}$
0.1227*	+0.964
.1092	.965
.1005	.973
.0503	.969
.0503	.970
.0402	.960

* HNO_3 was used instead of HClO_4 in this run.

It was found that increased acid concentration decreased the reproducibility of the potential, probably owing to attacking of the silver. Making the assumption that the values of the coefficient should be weighted inversely as the acid concentration, compute the weighted mean value for the temperature coefficient.

6. A determination of the atomic weight of chlorine by Hönigschmid and Chan ⁷ yielded these results:

35.4558	35.4565
.4563	.4562
.4553	.4568
.4574	.4558

Compute a and σ .

7. In their redetermination of the atomic weight of iodine, Hönigschmid and Striebel ⁸ obtained the following values for the ratio of silver iodide to silver chloride:

1.638085	1.638064	1.638078
058	085	080
070	070	067
100	106	083
078	080	066
061	069	

The arithmetic mean is 1.638076. Compute the standard deviation.

8. The data of Baxter and Hale (Problem 1) can be tabulated conveniently

⁶ *Ibid.*, **59**, 2271 (1937).

⁷ *Z. anorg. allgem. Chem.*, **163**, 315 (1927); see also Hönigschmid, *J. Am. Chem. Soc.*, **53**, 3012 (1931).

⁸ *Z. anorg. allgem. Chem.*, **208**, 53 (1932).

in a frequency table taking a class interval of 0.0010. To avoid confusion, the range of each class interval is given in the first column.

RANGE	ATOMIC WEIGHT	FREQUENCY
75- 84	12.0080	1
85- 94	12.0090	2
95-104	12.0100	11
105-114	12.0110	9
115-124	12.0120	3
125-134	12.0130	1

From this frequency table compute a and σ , and compare with (c) and (d) of Problem 1.

9. Plot the data of Problem 8 as a frequency graph, drawing a smooth curve through the points. How would you explain the departure of this distribution from the "normal" type shown in Fig. 9?

10. By means of vertical lines, mark on the graph prepared in Problem 9 the following values which were obtained in Problem 1: (a) the mode; (b) the median; (c) a ; (d) $a + \sigma$; and (e) $a - \sigma$.

CHAPTER IX

THEORY OF ERRORS

Truth is heavy, therefore few care to carry it.

—*The Talmud.*

In the preceding chapter we have examined the question of distribution fairly briefly. We have not succeeded, as yet, in defining distribution; we have merely described it. Perhaps if we can deduce a theoretical basis for distribution, and hence be able to express it mathematically, we shall approach a little nearer to an understanding.

It is a fact of observation that the values of a measured quantity are dispersed, the extent of dispersion depending upon the care observed and the skill of the experimenter. It is a further fact that the measured values fit into a pattern according to magnitude and frequency of occurrence. We must not be too hasty and say that all cases of the measurement of a quantity fit a distribution pattern. In measurements which consist in direct counting, perfect agreement in successive counts might be anticipated provided that the numbers are reasonably small. To be safe and also to make allowances for future improvements in experimental technique we shall state as an hypothesis: *distribution is an inherent property of most measurement, as now practiced.*

Distribution of Errors. If measurements fall into a distribution pattern, it follows that the errors of the measurements will likewise fit a distribution pattern of the same shape. It will be recalled that the variation of measurements is attributed to a large number of variables, either neglected or incompletely controlled, which in their free and practically independent operation result in deviations from the "true value." Inasmuch as deviations from the "true value" are the errors, the variations of the measurements and of the errors are both the results of one set of causative factors.

Measurements and their errors fit the same type of distribution pattern. Two examples will illustrate.

Example 1. If, for the following table of values of x and their frequencies f , we assume that the "true value" is 127, the column of d_{127} will be the "errors."

x	f	d_{127}
125	1	-2
126	3	-1
127	7	0
128	3	+1
129	1	+2

In Fig. 10, (a) is the frequency diagram of the measurements and (b) is the corresponding diagram for the errors. Except for a shift along the horizontal axis, the two curves are identical.

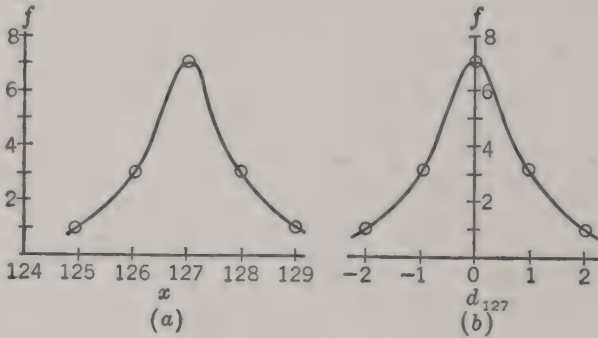


FIG. 10.

Example 2. In the calibration of a certain cryoscopic apparatus the following freezing points of water were obtained with the frequencies indicated:

$t^{\circ}\text{C}$	f	$t^{\circ}\text{C}$	f
0.003	2	-0.001	10
0.002	4	-0.002	3
0.001	9	-0.003	1
0.000	13	-0.004	1

Since by definition of the thermometer scale the freezing point of water is zero, then 0.000 is the "true value," and the measurements are the errors. It should not require graphs to show a correspondence of the frequency distributions of the measurements and of the errors in this instance.

Distribution Curves. The identity in type of distribution of measurements and of errors being recognized, it is apparent that,

if either curve can be fitted with a mathematical equation, by a suitable change of scale the equation will fit the other curve also. An equation of distribution has been derived based upon simple considerations from the mathematics of probability. A side excursion into the realm of mathematical probability will be the best approach.

PROBABILITY THEOREMS

Historical Statement. The mathematics of probabilities owes its origin to B. Pascal (1623–1662) and P. de Fermat (1601–1665), who jointly evolved the fundamental principles in an effort to solve the problem of how to divide the stakes of an unfinished game. Their discoveries were elaborated by C. Huygens (1629–1695), and further by James Bernoulli (1654–1705), whose posthumous book (1713) appears to be the first treatise on the subject. The great developer of the calculus of probabilities was P. Laplace (1749–1827), who in 1812 issued his *Théorie analytique des probabilités* which is still regarded as a classical work. Laplace called the theory of probabilities merely “common sense expressed in mathematical symbols.” It is simply a system for calculating results that might be deduced intuitively.

Mathematical Probability. If there is some doubt about the occurrence of an event, the calculus of probabilities can predict the chance that it will or will not happen, provided that the problem of occurrence does not reduce to a triviality. The question: “What is the chance that heads will appear when a coin is tossed?” presents a reasonable problem of probability since there is some doubt. Such a question as: “What is the probability that the number 2,152 has an integral square root?” is trivial because the square root can be taken and the answer is a statement of fact.

A quantitative answer to a proper problem in probability is obtained by means of the following reasoning:

- (1) If there are n equally likely ways in which an event can happen, and
- (2) if out of these n ways there are m ways in which a particularly designated event can occur; then
- (3) the probability p of the particular event occurring in a single trial is given by the ratio m/n .

That is,

$$p = \frac{m}{n} \quad (1)$$

This definition might be interpreted as presuming that the occurrence of the event had already been observed and tabulated. A more useful definition can be stated in terms of happenings and failures to happen. If an event can happen in k ways and fail to happen in l ways and all the $(k + l)$ ways are equally likely, the probability of happening is

$$p = \frac{k}{k + l} \quad (2)$$

and the probability of failing is

$$q = \frac{l}{k + l} \quad (3)$$

Example 1. A coin has two faces: heads and tails. If a coin is tossed there is one way in which heads can appear and one way in which heads can fail to appear, i.e., tails may appear. Therefore the probability of getting heads with a single toss of the coin is

$$p = \frac{1}{1 + 1} = \frac{1}{2}$$

Example 2. A die has six faces which are denoted 1, 2, 3, 4, 5, and 6. There is one way for 3 to show when the die is cast and five ways for the 3 to fail to appear. The probability of getting 3 with a single cast of the die is

$$p = \frac{1}{1 + 5} = \frac{1}{6}$$

Equally Likely Events. In the two examples just given, the reader has experienced no intuitive objection to considering heads and tails as equally likely to show. Neither does the intuition object to declaring each of the six faces of a die equally likely to appear. It is understood, of course, that these considerations could not apply to a coin which has heads on both sides or to a die that has been weighted. As is usual with intuitive conceptions, attempts to clarify with definitions and qualifications succeed chiefly in beclouding the issue of "equally likely." The first difficulty comes in decreeing whether by "equally likely" we mean "equal knowledge of likelihood" or "equal ignorance of likelihood." The two criteria should reduce to identity, but, in the prosecution

of these ideas to their logical extremes, further difficulties arise. For example, the question: "What is the probability that there are dogs on the dog-star?" would, on the basis of "equal ignorance," require the answer that "dogs" and "not-dogs" are equally likely.

A coin can be examined to determine whether the faces are alike or different and to judge of the symmetry or balance. We are totally ignorant of conditions on the "dog-star." Therefore, it is in judging of cases where the doubt approaches total ignorance that probability considerations must be applied with caution and skepticism toward the result.

Theorem 1. If the probability that an event will happen is p and the probability of its failing to happen is q , then

$$p + q = 1 \quad (4)$$

By definition from (2) and (3)

$$p = \frac{k}{k+l}; \quad q = \frac{l}{k+l}$$

Therefore

$$p + q = \frac{k}{k+l} + \frac{l}{k+l} = \frac{k+l}{k+l} = 1$$

Equation (4) may be restated as

$$p = 1 - q \quad (4a)$$

or

$$q = 1 - p \quad (4b)$$

Theorem 2. Unit probability corresponds to certainty. From the definition of probability, if there is no way in which an event can fail to happen, then

$$p = \frac{k}{k+0} = 1$$

Conversely, zero probability denotes certain failure of an event. From the above

$$q = \frac{0}{k+0} = 0$$

It is only in the realm of mathematics, however, that we can speak with assurance of unit or zero probability. In actuality the words "certain" and "impossible" are strong words which should be

applied with caution. In a coin toss, it is conceivable that the coin might land on edge. The statement of the probability of getting heads upon a single toss of a coin as exactly $\frac{1}{2}$ implies the condition "when the coin falls flat."

Complex Events. Many events can be resolved into several simpler events. Thus the event of obtaining an even number with one cast of a die may be resolved into three simpler events, namely: (1) obtaining 2, (2) obtaining 4, or (3) obtaining 6. It is obvious that, if any one of the three simple events occurs, neither of the other two can happen. Several events are denoted as **mutually exclusive** if one and only one of them can happen.

The resolution of complex events into simpler events gives rise to another classification according to whether the occurrence of one affects the occurrence of the others. If the occurrence of any one event has no effect on the occurrence of the others, the events are said to be **independent**. If, on the other hand, the occurrence of one event does affect the occurrence of the others, the events are said to be **related**. A simple example will illustrate. If a box contains 3 red and 5 green balls, the probability of a red ball being picked by a blindfolded person is $\frac{3}{8}$. If the ball is returned to the box and the trial repeated, the probability of getting a red ball is the same as before. The events are therefore independent. Suppose that after the first trial the ball were not returned. Upon a second trial the probability of drawing a red ball will be either $\frac{2}{7}$ or $\frac{3}{7}$, depending upon the color of the first ball that was removed. Here the events are related.

Theorem 3. The probability that any one of several mutually exclusive events will occur is the sum of their separate probabilities.

That is

$$p = p_1 + p_2 + p_3 + \cdots + p_n \quad (5)$$

This formula can be easily verified¹ for the occurrence of an even number upon one cast of a die. The probability of getting any particular number on the die is $\frac{1}{6}$. Since there are three even numbers, then

$$p = \frac{1}{6} + \frac{1}{6} + \frac{1}{6} = \frac{3}{6}$$

The same result can be obtained in another way. There are three

¹ For completely formal proofs of these theorems the reader should consult one of the textbooks on probability listed in the Appendix.

ways in which an even number can appear and three in which it can fail to appear, and all six ways are equally likely. From (2)

$$p = \frac{k}{k+l} = \frac{3}{3+3} = \frac{3}{6}$$

Theorem 4. The probability that several independent events will occur together is the product of their separate probabilities.

That is,

$$p = p_1 \cdot p_2 \cdot p_3 \cdots p_n \quad (6)$$

To “prove” this formula, let us examine an example of coin tossing. The simultaneous tossing of two coins may be considered the same as two separate tossings of one coin. What is the probability of getting two heads to appear when two coins are tossed? The probability of heads on one toss of one coin is $\frac{1}{2}$. From (6)

$$p = \frac{1}{2} \times \frac{1}{2} = \frac{1}{4}$$

Considering the same problem from another viewpoint we can obtain the identical answer. When two coins are tossed there is one way of getting two heads to appear and three ways for two heads to fail to appear. The possibilities are:

$$HH, HT, TH, TT$$

The probability according to (2) is

$$p = \frac{1}{1+3} = \frac{1}{4}$$

In the series of four possibilities just given we note the appearance of the two terms *HT* and *TH* which at first glance might appear to be only one term in two different guises. To clarify the matter we had better digress from our digression to review the meanings of permutations and combinations.

Permutations. We often desire to know in how many different ways a group of objects can be arranged either with the whole group or with parts of the group. Each possible arrangement is known as a permutation. For instance, how many numbers of three digits each can be made from the three digits 1, 2, and 3? The permutations of this group are 123, 132, 213, 231, 312, and 321. The two-digit numbers possible from the same three digits are 12, 21, 13, 31, 23, and 32.

In general, if there are n things, the number of permutations taking r at a time is given by

$${}_nP_r = n(n-1)(n-2)\cdots(n-r+1) \quad (7)$$

Equation (7) can be proved in this way: in an arrangement of r things, each of the n things can occupy the first position. There are therefore n ways of filling the first position. With any one of the n things occupying the first position, there are now only $(n-1)$ things left which can fill the second position. By similar reasoning there are $(n-2)$ ways of filling the third place. Observing that the number subtracted from n is 1 less than the number of the position, it follows that for the r th position there are $[n-(r-1)]$ or $(n-r+1)$ ways of filling the last position.

In the special case of n things taken n at a time, the last factor of the formula becomes

$$n - n + 1 = 1$$

Therefore

$${}_nP_n = n(n-1)(n-2)\cdots 1$$

The product of all the integers from n to 1 inclusive occurs frequently and it is called **factorial n** , denoted $n!$. Thus

$${}_nP_n = n! \quad (8)$$

Example 3. (a) The number of permutations of 3 numbers taken 3 at a time is

$${}_3P_3 = 3! = 3 \times 2 \times 1 = 6$$

(b) The number of permutations of 3 numbers taken 2 at a time is

$${}_3P_2 = 3(3-2+1) = 6$$

Both these answers agree with the tabulated possibilities at the first of this section.

Combinations. An arrangement of things taking no account of order is known as a combination. With the restriction of order removed there is only one possible combination of the integers 1, 2, and 3 taken three at a time. The six permutations of these integers, namely, 123, 132, 213, 231, 312, and 321, are all different arrangements of the one combination. When the 3 integers are taken 2 at a time, the possible combinations are 12, 13, and 23.

In general, the number of combinations of n things taken r at a time is given by

$${}_nC_r = \frac{n!}{(n-r)!r!} \quad (9)$$

For n things taken r at a time there are many more permutations possible than combinations. Consider one combination of r things and see how many arrangements of these r things are possible taking account of order. This amounts to asking how many permutations there are of r things taken r at a time. The answer according to (8) is

$${}_rP_r = r!$$

For every combination there are $r!$ permutations, or there are $r!$ times as many permutations as combinations.

$${}_nP_r = r!{}_nC_r \quad (10)$$

From which

$${}_nC_r = \frac{{}_nP_r}{r!}$$

From (7)

$${}_nC_r = \frac{n(n-1)(n-2)(n-3)\cdots(n-r+1)}{r!} \quad (11)$$

Multiplying both the numerator and denominator by $(n-r)!$

$${}_nC_r = \frac{n(n-1)(n-2)(n-3)\cdots(n-r+1)(n-r)!}{(n-r)!r!}$$

Since $(n-r)!$ is the product of all the integers from $(n-r)$ to 1 inclusive, and, since the series already in the numerator is the product of all the integers from n down to $(n-r+1)$, the second series fits in at the end of the first series to give the continuous product of all the integers from n to 1 inclusive. Hence

$${}_nC_r = \frac{n!}{(n-r)!r!}$$

To return for a moment to the series of possibilities for one toss of two coins discussed under Theorem 4, where we had the series

$$HH, HT, TH, TT$$

it might be argued that the two middle terms were permutations of letters instead of combinations of faces. If we denote one coin by capital letters and the second by small letters, the series becomes

$$Hh, Ht, Th, Tt$$

and it is apparent that the four possibilities are really equally likely.

Limitations of Probability Theorems. Once more the difference between the idealized conditions of mathematical postulation and the actual conditions of practice must be stressed. We have said the probability of getting heads when a coin is tossed is $\frac{1}{2}$. Every small boy who has "pitched pennies" knows better than to assume that, if heads comes once, tails will follow on the next throw. If a coin is tossed a large number of times, say 100, as many as 55 heads may result; if there are 1,000 tosses there may be 515 heads. As the number of throws is increased the ratio of heads to number of throws will approach $\frac{1}{2}$. The calculations of *a priori* probability represent only the limiting condition toward which the average approaches when the number of trials approaches infinity. For an observable event for which advance predictions are not possible the statistical average of occurrence may be used as an *a posteriori* estimate of the probability. Probabilities estimated in this way are called **empirical**. The combination of a man, a rifle, a target, and the remainder of the universe constitute so complex a system that it would be manifestly impossible to predict the probability of a bull's-eye on any particular shot. If, however, the result of tabulation of the man's scores over a long period of time were available and he had averaged 2 bull's-eyes out of each 10 shots, his empirically determined probability of hitting the bull's-eye on any single shot would be "about" $\frac{2}{10}$ or $\frac{1}{5}$.

COIN TOSSING

We shall first examine in detail all the possibilities of 1, 2, 3, and 4 tosses of a coin, the equivalents of a single toss of 1, 2, 3, or 4 coins, respectively, and see whether there is any regularity upon which we can generalize. If generalizations are possible, it should then be possible to predict the results of any number of coin tosses.

One Coin. With a single toss of one coin the possibilities are:

$$H, T$$

The probability of one heads is $\frac{1}{2}$, and the probability of one tails is $\frac{1}{2}$. The events are mutually exclusive, so by Theorem 3 the probability that one event will occur is

$$P = \frac{1}{2} + \frac{1}{2}$$

Two Coins. For a single toss of two coins, the possibilities are

2 heads	1 heads	0 heads
HH	HT	TT
	TH	

The probability of HH is the product of the separate probabilities of the two independent events (Theorem 4) or $(\frac{1}{2})(\frac{1}{2}) = (\frac{1}{2})^2$. The probability of HT is the product of the probability of heads ($\frac{1}{2}$) and the probability of tails ($\frac{1}{2}$) or $(\frac{1}{2})(\frac{1}{2})$. But since there is another possibility, for which the probability is $(\frac{1}{2})(\frac{1}{2})$, which is just as likely as HT , and since the events HT and TH are mutually exclusive, the probability that one heads and one tails will appear is the sum of the separate probabilities, or

$$p = 2(\frac{1}{2})(\frac{1}{2})$$

We could have obtained the same result by considering how many combinations of heads and tails could give exactly one heads in each combination. From (9)

$${}_2C_1 = \frac{2!}{1!1!} = \frac{2 \times 1}{1 \times 1} = 2$$

By reasoning similar to the case of HH , the probability of TT is equal to $(\frac{1}{2})^2$.

Now, because the possibilities two heads, one heads, and no heads are mutually exclusive, the probability of one possibility is the sum of the separate probabilities. That is,

$$P = (\frac{1}{2})^2 + 2(\frac{1}{2})(\frac{1}{2}) + (\frac{1}{2})^2$$

Three Coins. The possibilities in one toss of three coins are

3 heads	2 heads	1 heads	0 heads
HHH	HHT	HTT	TTT
	HTH	THT	
	THH	TTH	

The probability of three heads is $(\frac{1}{2})(\frac{1}{2})(\frac{1}{2}) = (\frac{1}{2})^3$.

The probability of two heads is $(\frac{1}{2})(\frac{1}{2}) = (\frac{1}{2})^2$; the probability of one tails is $\frac{1}{2}$; hence the probability of *HHT* is $(\frac{1}{2})^2 (\frac{1}{2})$. The number of equivalent combinations is

$${}_3C_2 = \frac{3!}{1!2!} = \frac{3 \times 2 \times 1}{1 \times 2 \times 1} = 3$$

The probability of two heads when three coins are tossed is $3(\frac{1}{2})^2 (\frac{1}{2})$.

By similar reasoning the probability of one heads when three coins are tossed is

$${}_3C_1 \left(\frac{1}{2}\right) \left(\frac{1}{2}\right)^2 = \frac{3!}{2!1!} \left(\frac{1}{2}\right) \left(\frac{1}{2}\right)^2 = 3 \left(\frac{1}{2}\right) \left(\frac{1}{2}\right)^2$$

Continuing with reasoning identical to the possibility of three heads; for zero heads the probability is $(\frac{1}{2})^3$.

As a final result, the probability of one possibility, namely, 3 heads, 2 heads, 1 heads, or 0 heads, is

$$P = (\frac{1}{2})^3 + 3(\frac{1}{2})^2(\frac{1}{2}) + 3(\frac{1}{2})(\frac{1}{2})^2 + (\frac{1}{2})^3.$$

Four Coins. Without any more detailed consideration than to tabulate the possibilities:

4 heads	3 heads	2 heads	1 heads	0 heads
<i>HHHH</i>	<i>HHHT</i>	<i>HHTT</i>	<i>HTTT</i>	<i>TTTT</i>
	<i>HHTH</i>	<i>HTHT</i>	<i>TTTH</i>	
	<i>HTHH</i>	<i>HTTH</i>	<i>THTT</i>	
	<i>THHH</i>	<i>THTH</i>	<i>TTHT</i>	
		<i>TTHH</i>		
		<i>THHT</i>		

we shall state that for a single toss of four coins the probability that one of the five possible numbers of heads will be obtained is

$$P = (\frac{1}{2})^4 + 4(\frac{1}{2})^3(\frac{1}{2}) + 6(\frac{1}{2})^2(\frac{1}{2})^2 + 4(\frac{1}{2})(\frac{1}{2})^3 + (\frac{1}{2})^4$$

Binomial Expansion. Any textbook of college algebra gives the binomial theorem for the expression of a sum raised to a power, by means of a series of $(n + 1)$ terms, where n is the power. The

binomial theorem was discovered by Sir Isaac Newton and first communicated to Leibnitz in 1676.² In general form,

$$(a + b)^n = a^n + {}_nC_1 a^{n-1} b + {}_nC_2 a^{n-2} b^2 + \cdots + {}_nC_{n-1} a b^{n-1} + b^n \quad (12)$$

Let us carry out the expansion of the sum $(\frac{1}{2} + \frac{1}{2})$ to several powers:

$$(1) \quad (\frac{1}{2} + \frac{1}{2})^1 = \frac{1}{2} + \frac{1}{2}$$

$$(2) \quad (\frac{1}{2} + \frac{1}{2})^2 = (\frac{1}{2})^2 + 2(\frac{1}{2})(\frac{1}{2}) + (\frac{1}{2})^2$$

$$(3) \quad (\frac{1}{2} + \frac{1}{2})^3 = (\frac{1}{2})^3 + 3(\frac{1}{2})^2(\frac{1}{2}) + 3(\frac{1}{2})(\frac{1}{2})^2 + (\frac{1}{2})^3$$

$$(4) \quad (\frac{1}{2} + \frac{1}{2})^4 = (\frac{1}{2})^4 + 4(\frac{1}{2})^3(\frac{1}{2}) + 6(\frac{1}{2})^2(\frac{1}{2})^2 + 4(\frac{1}{2})(\frac{1}{2})^3 + (\frac{1}{2})^4$$

For comparison we shall summarize the results of our consideration of the tossing of 1, 2, 3, or 4 coins.

$$(1 \text{ coin}) \quad P = \frac{1}{2} + \frac{1}{2}$$

$$(2 \text{ coins}) \quad P = (\frac{1}{2})^2 + 2(\frac{1}{2})(\frac{1}{2}) + (\frac{1}{2})^2$$

$$(3 \text{ coins}) \quad P = (\frac{1}{2})^3 + 3(\frac{1}{2})^2(\frac{1}{2}) + 3(\frac{1}{2})(\frac{1}{2})^2 + (\frac{1}{2})^3$$

$$(4 \text{ coins}) \quad P = (\frac{1}{2})^4 + 4(\frac{1}{2})^3(\frac{1}{2}) + 6(\frac{1}{2})^2(\frac{1}{2})^2 + 4(\frac{1}{2})(\frac{1}{2})^3 + (\frac{1}{2})^4$$

We immediately perceive that, in each of the four cases, the over-all probability P can be represented by

$$P = (\frac{1}{2} + \frac{1}{2})^n$$

where n is the number of coins tossed at once or the number of tosses of one coin. We might generalize further and say that, since the probability of heads (p) is $\frac{1}{2}$ and the probability of tails ($1 - p = q$) is $\frac{1}{2}$, then

$$P = (p + q)^n \quad (13)$$

These generalizations have been somewhat hasty, and we have as yet no grounds for calling (13) a general formula. Suffice it to say that it has been well established as generally true, and it is one of the most useful formulas in statistics. Interested readers can find generalized proofs and complete discussions of "binomial distribution" in almost any one of the statistics books listed in the

² Ball: *History of Mathematics*, Second Edition, p. 328, Macmillan, London, 1893.

Appendix. The probabilities of obtaining $n, n - 1, \dots, 1, 0$ of any one face of a die when n dice are thrown are given by the successive terms of the expansion of $(\frac{1}{6} + \frac{5}{6})^n$.

Frequency Distribution in Coin Tossing. Let us examine once more the expansion for the tossing of four coins. The series is:

$$(\frac{1}{2})^4 + 4(\frac{1}{2})^3(\frac{1}{2}) + 6(\frac{1}{2})^2(\frac{1}{2})^2 + 4(\frac{1}{2})(\frac{1}{2})^3 + (\frac{1}{2})^4$$

The first term gives the probability of 4 heads, which equals $\frac{1}{16}$. Recalling the fundamental "definition" of probability [equation (1)], we see that 4 heads occurs 1 time out of 16. If 16 is the total number of observations, 1 is the frequency of 4 heads if we consider 4 heads as the observed value. The first term $\frac{1}{16}$ gives not the frequency but a number which is proportional to the frequency; we might call it the "frequency coefficient." Likewise, for the second term, $\frac{4}{16}$ is the frequency coefficient of 3 heads. Arguing again toward generalizations, the successive terms of an expansion of $(p + q)^n$ give the frequency coefficients of $n, n - 1, n - 2, \dots, 1, 0$, respectively, of the event being observed—heads of the coin in this instance.

For purposes of computation, it will be permissible to treat the frequency coefficient as if it were the frequency itself since the coefficient is directly proportional to the frequency. Let us examine a frequency table for 1 toss of 4 coins. The frequency coefficient will hereafter be denoted by F to distinguish it from the frequency (f). Let x denote the number of heads, and take 2 as x' .

x	F	$d_{x'}$	$Fd_{x'}$	$Fd_{x'}^2$
4	1/16	2	2/16	4/16
3	4/16	1	4/16	4/16
2	6/16	0	0	0
1	4/16	-1	-4/16	4/16
0	1/16	-2	-2/16	4/16
$\Sigma F = 1$			$\Sigma Fd_{x'} = 0$	$\Sigma Fd_{x'}^2 = 1$

$$a = 2 - \frac{0}{1} = 2$$

$$\sigma = \pm \sqrt{\frac{1}{1} - \left(\frac{0}{1}\right)^2} = \pm 1$$

We should have expected to find the arithmetic mean for this distribution to be 2.

Average and Dispersion in a Binomial Distribution. For the expansion of $(p + q)^n$, in which p is the probability of heads and q the probability of tails on a single toss and n is the number of coins tossed, we shall construct a frequency table just as we did for the tossing of 4 coins. The expanded series is:

$$(p + q)^n = p^n + {}_nC_1 p^{n-1} q + {}_nC_2 p^{n-2} q^2 + \cdots + {}_nC_{n-2} p^2 q^{n-2} \\ + {}_nC_{n-1} p q^{n-1} + q^n$$

Solving for the various values of ${}_nC_r$, we obtain:

$${}_nC_1 = \frac{n!}{(n-1)!1!} = n$$

$${}_nC_2 = \frac{n!}{(n-2)!2!} = \frac{n(n-1)}{2!}$$

$$\begin{array}{ccc} \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot \end{array}$$

$${}_nC_{n-2} = \frac{n!}{2!(n-2)!} = \frac{n(n-1)}{2!}$$

$${}_nC_{n-1} = \frac{n!}{1!(n-1)!} = n$$

The series becomes:

$$(p + q)^n = p^n + np^{n-1}q + \frac{n(n-1)}{2!} p^{n-2}q^2 + \cdots \\ + \frac{n(n-1)}{2!} p^2q^{n-2} + npq^{n-1} + q^n$$

Let x be the number of heads, and let $x' = 0$.

x	F	d_x	Fd_x	Fd_x^2
n	p^n	n	np^n	n^2p^n
$n-1$	$np^{n-1}q$	$n-1$	$n(n-1)p^{n-1}q$	$n(n-1)^2p^{n-1}q$
$n-2$	$\frac{n(n-1)}{2!}p^{n-2}q^2$	$n-2$	$\frac{n(n-1)(n-2)}{2!}p^{n-2}q^2$	$\frac{n(n-1)(n-2)^2}{2!}p^{n-2}q^2$
$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$
2	$\frac{n(n-1)}{2!}p^2q^{n-2}$	2	$\frac{2n(n-1)}{2!}p^2q^{n-2}$	$\frac{4n(n-1)}{2!}p^2q^{n-2}$
1	npq^{n-1}	1	npq^{n-1}	npq^{n-1}
0	q^n	0	0	0

To obtain the arithmetic mean,

$$a = x' + \frac{\Sigma Fd_x}{\Sigma F}$$

But

$$x' = 0$$

And

$$\Sigma F = (p + q)^n = 1$$

Therefore,

$$a = \Sigma Fd_x$$

$$a = np^n + n(n-1)p^{n-1}q + \frac{n(n-1)(n-2)}{2!}p^{n-2}q^2 + \dots$$

$$+ \frac{2n(n-1)}{2!}p^2q^{n-2} + npq^{n-1}$$

Factoring out np ,

$$a = np \left[p^{n-1} + (n-1)p^{n-2}q + \frac{(n-1)(n-2)}{2!}p^{n-3}q^2 + \dots \right.$$

$$\left. + (n-1)pq^{n-2} + q^{n-1} \right]$$

The terms in the bracket are the terms of the expansion of

$$(p + q)^{n-1} = 1$$

Therefore

$$a = np \quad (14)$$

The root mean square deviation from an arbitrary value x' is given by

$$s^2 = \frac{\Sigma F d_{x'}^2}{\Sigma F}$$

Since $\Sigma F = 1$

$$s^2 = \Sigma F d_{x'}^2$$

$$\begin{aligned} s^2 = n^2 p^n + n(n-1)^2 p^{n-1} q + \frac{n(n-1)(n-2)^2}{2!} p^{n-2} q^2 + \dots \\ + \frac{4n(n-1)}{2!} p^2 q^{n-2} + npq^{n-1} \end{aligned}$$

Factoring out np ,

$$\begin{aligned} s^2 = np \left[np^{n-1} + (n-1)^2 p^{n-2} q + \frac{(n-1)(n-2)^2}{2!} p^{n-3} q^2 + \dots \right. \\ \left. + \frac{4(n-1)}{2!} p q^{n-2} + q^{n-1} \right] \end{aligned}$$

We shall need to split some of the terms of the series as indicated by the following steps.

$$(1) \quad np^{n-1} = [(n-1) + 1] p^{n-1} = (n-1) p^{n-1} + p^{n-1}$$

$$\begin{aligned} (2) \quad (n-1)^2 p^{n-2} q &= [(n-2) + 1] (n-1) p^{n-2} q \\ &= (n-1)(n-2) p^{n-2} q + (n-1) p^{n-2} q \end{aligned}$$

$$\begin{aligned} (3) \quad \frac{(n-1)(n-2)^2}{2!} p^{n-3} q^2 &= [(n-3) + 1] \frac{(n-1)(n-2)}{2!} p^{n-3} q^2 \\ &= \frac{(n-1)(n-2)(n-3)}{2!} p^{n-3} q^2 \\ &\quad + \frac{(n-1)(n-2)}{2!} p^{n-3} q^2 \end{aligned}$$

$$\begin{aligned} (4) \quad \frac{4(n-1)}{2!} p q^{n-2} &= 2(n-1) p q^{n-2} \\ &= (n-1) p q^{n-2} + (n-1) p q^{n-2} \end{aligned}$$

Substituting for the original terms of the series, the expression becomes

$$s^2 = np \left[p^{n-1} + (n-1)p^{n-1} + (n-1)p^{n-2}q + (n-1)(n-2)p^{n-2}q \right. \\ \left. + \frac{(n-1)(n-2)}{2!} p^{n-3}q^2 + \frac{(n-1)(n-2)(n-3)}{2!} p^{n-3}q^2 \right. \\ \left. + \cdots + (n-1)pq^{n-2} + (n-1)pq^{n-2} + q^{n-1} \right]$$

Collecting the odd-numbered terms in one series and the even-numbered terms in another series,

$$s^2 = np \left[p^{n-1} + (n-1)p^{n-2}q + \frac{(n-1)(n-2)}{2!} p^{n-3}q^2 + \cdots \right. \\ \left. + (n-1)pq^{n-2} + q^{n-1} \right] + np \left[(n-1)p^{n-1} + (n-1)(n-2)p^{n-2}q \right. \\ \left. + \frac{(n-1)(n-2)(n-3)}{2!} p^{n-3}q^2 + \cdots + (n-1)pq^{n-2} \right]$$

The terms of the first series are the terms of the expansion of $(p + q)^{n-1}$, which is equal to 1. The second series has a common factor $(n-1)p$ which can be factored.

$$s^2 = np + np \left[(n-1)p \left\{ p^{n-2} + (n-2)p^{n-3}q \right. \right. \\ \left. \left. + \frac{(n-2)(n-3)}{2!} p^{n-4}q^2 + \cdots + q^{n-2} \right\} \right]$$

The terms of the remaining series are the terms of the expansion of $(p + q)^{n-2}$, which is equal to 1. Therefore

$$s^2 = np + np[(n-1)p] = np + n^2p^2 - np^2$$

By definition

$$\sigma^2 = s^2 - (a - x')^2$$

Since $x' = 0$ and $a = np$, hence

$$\sigma^2 = np + n^2p^2 - np^2 - n^2p^2$$

$$\sigma^2 = np(1 - p)$$

Because $(1 - p) = q$

$$\begin{aligned}\sigma^2 &= npq \\ \sigma &= \sqrt{npq}\end{aligned}\tag{15}$$

Although the lengthy derivation required to obtain (14) and (15) looks like "rabbits being pulled from a hat," the mathematics is perfectly sound and the fundamental importance of equations (14) and (15) amply justifies the tedium. Let us check the values of a and σ for the tossing of 4 coins (p. 154). For a single toss of 4 coins

$$p = \frac{1}{2}, \quad q = \frac{1}{2}, \quad n = 4$$

$$a = np = 4 \times \frac{1}{2} = 2$$

$$\sigma = \pm \sqrt{npq} = \pm \sqrt{4 \times \frac{1}{2} \times \frac{1}{2}} = \pm 1$$

The results agree perfectly with those obtained on page 154.

CURVE OF ERROR

Having obtained the expression $(p + q)^n$, the binomial expansion of which gives the distribution of the frequency coefficients for the occurrence of various measured values of a certain quantity (heads of a coin, faces of a die, etc.), and being able also to express the dispersion of those values, we should be now in a position to obtain a generalized expression for normal distribution of frequency coefficients.

In Chapter VI it was postulated that errors of measurement are the resultants of a great many neglected and imperfectly controlled variables operating practically independently of one another. This is equivalent to saying each variable is responsible for an "error" of secondary magnitude which may be either positive or negative and of varying size. To simulate the conditions of a measurement we might imagine a very large number of coins of various denominations and think of heads as positive and tails as negative. Then, considering each of the possible combinations of heads and tails with the varying possible algebraic sums, we can get a glimpse at the theoretical background of errors as now understood. In the derivation³ which follows, we have no way of postulating coins of

³ This derivation is based on the treatment of Gale and Watkeys: *Elementary Functions*, Chapter IX, Henry Holt, New York, 1920.

different denominations. The best approximation we can get is the assumption of coins all alike.

A plot of the values of the terms of the expansion of

$$(p + q)^n$$

when n is a small number, results in a frequency polygon, not a continuous curve. Consider the expansion of

$$\left(\frac{1}{2} + \frac{1}{2}\right)^{2n}$$

We choose $2n$ power for convenience. The number of terms of the expansion is always one greater than the power; hence $2n + 1$ will always be an odd number and insure a middle term.

From (15)

$$\sigma = \sqrt{npq}$$

For the above expansion

$$\sigma = \sqrt{2n \cdot \frac{1}{2} \cdot \frac{1}{2}} = \sqrt{\frac{n}{2} \cdot 1 \cdot 1}$$

or

$$2\sigma^2 = n(1)^2$$

Since n is going to be a very large number, if the form of the binomial expansion (12) is recalled, it will be apparent that the magnitudes of the frequency coefficients will become very small because the sum of all the terms is 1. In order to prevent the curve from becoming very flat, we can substitute for the interval of unity in the values of x a generalized interval Δx . Since σ is a constant expressing the dispersion, we can take

$$2\sigma^2 = n(\Delta x)^2 \quad (16)$$

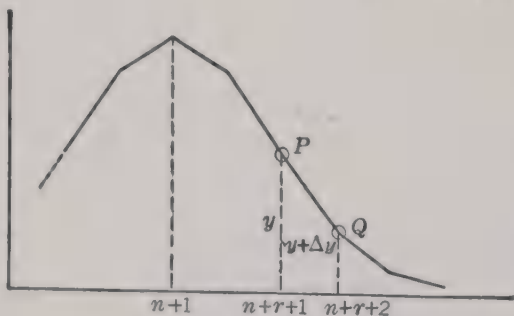


FIG. 11.

and let Δx decrease in such a way that, as n increases, equality of the right- and left-hand sides of (16) will be maintained.

Taking n as some finite value, let us consider two terms of the expanded series. We shall let P and Q be two successive vertices of the frequency

polygon determined by the r th and the $(r + 1)$ th terms from the center term of the expansion, the $(n + 1)$ th term. Fig. 11 represents this. From equation (12)

At P :

$$y = {}_{2n}C_{n+r} \left(\frac{1}{2}\right)^{2n}$$

At Q :

$$y + \Delta y = {}_{2n}C_{n+r+1} \left(\frac{1}{2}\right)^{2n}$$

From the definition of ${}_nC_r$,

$$y = \left(\frac{1}{2}\right)^{2n} \frac{(2n)!}{(n-r)!(n+r)!}$$

$$y + \Delta y = \left(\frac{1}{2}\right)^{2n} \frac{(2n)!}{(n-r-1)!(n+r+1)!}$$

By division,

$$\frac{y + \Delta y}{y} = \frac{(n-r)!(n+r)!}{(n-r-1)!(n+r+1)!} = \frac{n-r}{n+r+1}$$

Then

$$\Delta y = \frac{n-r}{n+r+1} y - y = \left[\frac{n-r}{n+r+1} - 1 \right] y$$

$$\Delta y = \left[\frac{n-r-(n+r+1)}{n+r+1} \right] y = - \frac{2r+1}{n+r+1} y$$

Divide each side by Δx

$$\frac{\Delta y}{\Delta x} = - \frac{2r+1}{n+r+1} \cdot \frac{y}{\Delta x}$$

Now multiply the right-hand side by $\frac{\Delta x}{\Delta x}$,

$$\frac{\Delta y}{\Delta x} = - \frac{2r\Delta x + \Delta x}{n(\Delta x)^2 + r(\Delta x)^2 + (\Delta x)^2} \cdot y$$

The value of x corresponding to P is

$$x = r\Delta x$$

and

$$2\sigma^2 = n(\Delta x)^2$$

Making these substitutions,

$$\frac{\Delta y}{\Delta x} = - \frac{2x + \Delta x}{2\sigma^2 + x\Delta x + (\Delta x)^2} \cdot y$$

As n becomes very large, Δx approaches zero, hence at the limit

$$\frac{dy}{dx} = - \frac{2xy}{2\sigma^2} = - \frac{xy}{\sigma^2}$$

Rearranging

$$\frac{dy}{y} = - \frac{x dx}{\sigma^2}$$

Integration gives

$$\ln y = - \frac{x^2}{2\sigma^2} + \ln C$$

or

$$y = Ce^{-\frac{x^2}{2\sigma^2}} \quad (17)$$

Properties of the Curve of Error. In order to see what the curve of error is like, we shall need to examine some of the mathematical properties of equation (17), which belongs to the class known as exponential equations. First of all, because x in the exponent is squared, substituting either $x = a$ or $x = -a$ will give the same value of y . Hence the curve expressed by equation (17) is symmetrical about the ordinate at $x = 0$. The ordinate at $x = 0$ will be

$$y = Ce^0 = C$$

As the value of x increases, the value of the factor $(e^{-\frac{x^2}{2\sigma^2}})$ decreases. As x approaches infinity, y approaches zero, and it is only when $x = \infty$ that y becomes zero. That is,

$$e^{-\infty} = 0$$

Therefore the points at which the curve intersects the x axis are at $x = \infty$ and $x = -\infty$.

The general shape of a distribution curve shows curvature downward at the middle and upward at the two extremes. It will be interesting to locate the positions of change of curvature, that is, the points of inflexion. The point of inflexion is the point at

which the rate of change of slope is zero, in other words, the point at which the second derivative is zero. For

$$y = Ce^{-\frac{x^2}{2\sigma^2}}$$

$$\frac{dy}{dx} = Ce^{-\frac{x^2}{2\sigma^2}} \left(-\frac{x}{\sigma^2} \right)$$

$$\frac{dy}{dx} = -\frac{Cx}{\sigma^2} e^{-\frac{x^2}{2\sigma^2}}$$

The second derivative is

$$\frac{d^2y}{dx^2} = -\frac{C}{\sigma^2} \left[xe^{-\frac{x^2}{2\sigma^2}} \left(-\frac{x}{\sigma^2} \right) + e^{-\frac{x^2}{2\sigma^2}} \right]$$

or

$$\frac{d^2y}{dx^2} = \frac{Ce^{-\frac{x^2}{2\sigma^2}}}{\sigma^2} \left(\frac{x^2}{\sigma^2} - 1 \right)$$

Setting the second derivative equal to zero,

$$\frac{Ce^{-\frac{x^2}{2\sigma^2}}}{\sigma^2} \left(\frac{x^2}{\sigma^2} - 1 \right) = 0$$

Since there are only two possible values of x for which the exponential factor could become zero, namely, $x = \infty$ and $x = -\infty$, then it must be true that

$$\frac{x^2}{\sigma^2} - 1 = 0$$

from which

$$x = \pm \sigma \quad (18)$$

The curvature changes at the points on the curve corresponding to $x = \sigma$ and $x = -\sigma$. The constant σ which expresses the dispersion is seen to determine these unique points of the curve of error. (See again the graph you prepared in Problem 10, Set XII.)

According to the original assumptions of the derivation of equation (17), an error of magnitude x has its frequency coefficient given by the corresponding ordinate y . This amounts to the statement that an error x , which falls in the infinitesimal class interval dx , will have its probability given by the area of a rectangle of height y and width dx . Or

$$p_1 = y_1 dx$$

Inasmuch as the occurrence of an error of some size is certain, and certainty represents unit probability, then, by the theorem of mutually exclusive events, the sum of all the separate probabilities must equal unity. That is,

$$\sum_{-\infty}^{\infty} y \, dx = 1$$

By the integral calculus

$$\sum_{-\infty}^{\infty} y \, dx = \int_{-\infty}^{\infty} y \, dx = 1$$

Further

$$\int_{-\infty}^{\infty} y \, dx = C \int_{-\infty}^{\infty} e^{-\frac{x^2}{2\sigma^2}} \, dx = 1$$

Almost any table of definite integrals will reveal that

$$\int_{-\infty}^{\infty} e^{-t^2} \, dt = \sqrt{\pi}$$

In our exponential factor

$$t = \frac{x}{\sqrt{2} \sigma}$$

and

$$dt = \frac{dx}{\sqrt{2} \sigma}$$

hence

$$\int_{-\infty}^{\infty} e^{-\frac{x^2}{2\sigma^2}} \frac{dx}{\sqrt{2} \sigma} = \sqrt{\pi}$$

It follows that

$$\int_{-\infty}^{\infty} e^{-\frac{x^2}{2\sigma^2}} \, dx = \sqrt{2\pi} \sigma$$

But

$$C \sqrt{2\pi} \sigma = 1$$

Therefore

$$C = \frac{1}{\sqrt{2\pi} \sigma} \quad (19)$$

The value for C being substituted in equation (17), the result is

$$y = \frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{x^2}{2\sigma^2}} \quad (20)$$

We note from equation (19) that the value of C which gives the maximum ordinate of the curve (at $x = 0$) is inversely proportional to σ . Recalling also that σ determines the location of the points of inflexion [equation (18)] of the curve, it will be interesting to see how the shape of the curve varies when σ is large or small. In Fig. 12, curve (1) shows a plot of the function given by (20)

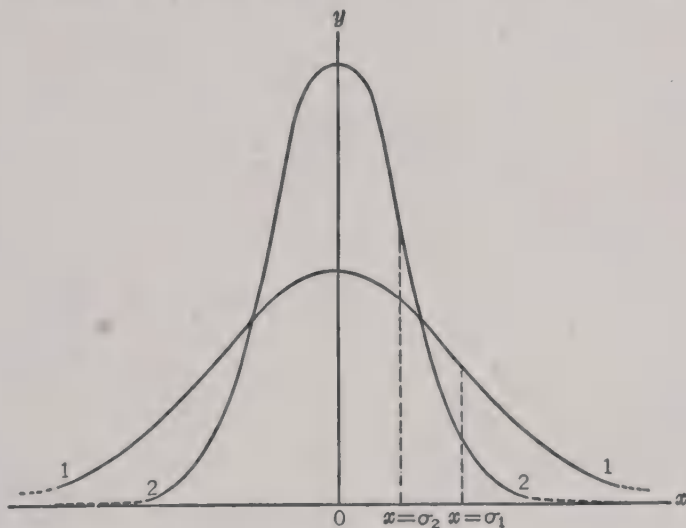


FIG. 12.

when σ has a large value, and curve (2) is for a small value of σ . The areas under the two curves are the same.

The equation for the curve of error in its final form [equation (20)] is known as the "Gaussian law of error." The equation was first stated by Gauss in 1795, but only as an empirical theorem. It remained for Laplace to work out a proof of the theorem, the publication of which occurred in 1812.⁴ Earlier efforts to deduce a "law" of error had been made by R. Cotes (1682–1716), by T. Simpson (1710–1761), and by J. Lagrange (1736–1813). All these early efforts were directed toward the fitting of an empirical curve to the "errors" of observations. Laplace appears to have been the first to recognize fully the probability aspect of errors.

⁴ Ball, *op. cit.*, p. 428.

Principle of Least Squares. The Gaussian "law" of error is sometimes called the "principle of least squares." The confusion arises from the fact that the "principle" is *implicitly* stated in the "law." A. M. Legendre (1752-1833) worked out the principle of least squares in 1806 independently of Gauss. The principle also awaited a formal proof by Laplace in 1812.

In essence, the principle of least squares is this: if for a group of measurements $(x_1, x_2, x_3, \dots, x_n)$, supposed to have been measured with uniform exactness, the most probable value of the quantity is assumed to be z , then the deviations of those measurements are given by $(x_1 - z)$, $(x_2 - z)$, $\dots$ $(x_n - z)$, respectively. Assuming for the moment that the deviations are the errors, the probability of the deviation of x_1 from z is, by equation (17),

$$y_1 = Ce^{-\frac{(x_1 - z)^2}{2\sigma^2}}$$

And likewise

$$y_2 = Ce^{-\frac{(x_2 - z)^2}{2\sigma^2}}; \dots; y_n = Ce^{-\frac{(x_n - z)^2}{2\sigma^2}}$$

These errors are independent; hence the probability of obtaining all the errors is the product of the separate probabilities or

$$P = Ce^{-\frac{(x_1 - z)^2}{2\sigma^2}} \cdot Ce^{-\frac{(x_2 - z)^2}{2\sigma^2}} \dots Ce^{-\frac{(x_n - z)^2}{2\sigma^2}}$$

Collecting exponents by the law of summation of exponents,

$$P = C^n e^{-\frac{1}{2\sigma^2}[(x_1 - z)^2 + (x_2 - z)^2 + \dots + (x_n - z)^2]}$$

Recalling that the value of an exponential factor such as e^{-t} becomes greater as t becomes smaller, it is apparent from the above equation that the probability (P) will be maximum when the portion of the exponent in the bracket is a minimum, or

$$(x_1 - z)^2 + (x_2 - z)^2 + \dots + (x_n - z)^2 = Q$$

must be a minimum.

The conditions for a minimum according to the differential calculus are

$$\frac{dQ}{dz} = 0 \quad \text{and} \quad \frac{d^2Q}{dz^2} = +$$

In words, at a minimum point of a curve the slope is zero (tangent at the point is horizontal) and the rate of change of slope away

from the minimum must be a positive rate. Let us test the function Q .

$$\frac{dQ}{dz} = -2(x_1 - z) - 2(x_2 - z) - \cdots - 2(x_n - z) = 0$$

and

$$\frac{d^2Q}{dz^2} = 2 + 2 + \cdots + 2 = 2n$$

The second derivative is positive, hence there is a minimum.

Returning to the equation for the first derivative, the condition of setting the expression equal to zero fixes the critical value for the minimum.

$$2(x_1 - z) + 2(x_2 - z) + \cdots + 2(x_n - z) = 0$$

which after factoring out the 2 becomes

$$nz = \Sigma x$$

Then

$$z = \frac{\Sigma x}{n} = a \quad (21)$$

The conclusions are therefore: (1) the most probable value to represent a set of measurements which have been made with equal exactness is the arithmetic mean; and (2) the sum of the squares of the deviations from the arithmetic mean is smaller than the sum of the squares of the deviations from any other value.

The Probability Integral. The probability integral, sometimes called the "error function" or the "error integral," may be found in other textbooks and in handbooks in one of these three forms:

$$\Phi(t) = \frac{1}{\sqrt{\pi}} \int_{-t}^t e^{-t^2} dt \quad (22)$$

$$\Phi(t) = \frac{2}{\sqrt{\pi}} \int_0^t e^{-t^2} dt \quad (23)$$

$$F(x) = \frac{2}{\sqrt{\pi}} \int_0^{hx} e^{-h^2 x^2} h dx \quad (24)$$

All these forms are equivalent. (22) and (23) are equivalent because the curve of error is symmetrical about the ordinate at

($x = 0$). The area under the curve between $-t$ and 0 is the same as the area between 0 and $+t$, hence the area between $-t$ and $+t$ is twice that between 0 and t . The third form is the same function as (22) and (23) but with the argument hx instead of t ; h is the so-called precision constant used by several writers on the subject of errors. The relation between h and σ is given by

$$h = \frac{1}{\sqrt{2}\sigma} \quad (25)$$

We shall make use only of the first form (22).

The significance of the probability integral may be interpreted in the following fashion. It has already been shown that the probability of an error of any magnitude between $-\infty$ and $+\infty$ is given by

$$P = \frac{1}{\sqrt{\pi}} \int_{-\infty}^{\infty} e^{-\frac{x^2}{2\sigma^2}} \frac{dx}{\sqrt{2}\sigma} = 1$$

By the same reasoning, each of the possible errors between the finite limits $-x_1$ and x_1 is a mutually exclusive event; hence the probability that an error will lie between these limits is given by

$$P = \frac{1}{\sqrt{\pi}} \int_{-\frac{x_1}{\sqrt{2}\sigma}}^{\frac{x_1}{\sqrt{2}\sigma}} e^{-\frac{x^2}{2\sigma^2}} \frac{dx}{\sqrt{2}\sigma}$$

Since

$$\frac{x}{\sqrt{2}\sigma} = t$$

it will be simpler to use the formula (22) with t as the argument,

$$P = \frac{1}{\sqrt{\pi}} \int_{-t_1}^{t_1} e^{-t^2} dt$$

The evaluation of this integral with finite limits to any desired number of significant figures can be accomplished by first expanding the exponential factor into an infinite series by means of Taylor's theorem.

$$e^{-t^2} = 1 - \frac{t^2}{1!} + \frac{t^4}{2!} - \frac{t^6}{3!} + \dots$$

The integral becomes

$$\Phi(t) = \frac{1}{\sqrt{\pi}} \int_{-t_1}^{t_1} \left(1 - \frac{t^2}{1!} + \frac{t^4}{2!} - \frac{t^6}{3!} + \dots \right) dt$$

Upon integration of successive terms of the series the expression becomes

$$\Phi(t) = \frac{1}{\sqrt{\pi}} \left[t - \frac{t^3}{1! \cdot 3} + \frac{t^5}{2! \cdot 5} - \frac{t^7}{3! \cdot 7} + \dots \right]_{-t_1}^{t_1}$$

Applying the limits, the equation simplifies to

$$\Phi(t_1) = \frac{2}{\sqrt{\pi}} \left[t_1 - \frac{t_1^3}{1! \cdot 3} + \frac{t_1^5}{2! \cdot 5} - \frac{t_1^7}{3! \cdot 7} + \dots \right] \quad (26)$$

The values of $\Phi(t)$ for various successive values of t have been computed and assembled for reference in tabular form. Table II in the Appendix gives a compilation of values for the probability integral.

Measure of Precision: The Probable Error. In Chapter VI the significance of precision as expressing the degree of concordance or consistency within a set of observed values was discussed. In a certain sense, the standard deviation, expressing as it does the dispersion, gives a measure of consistency. For statistical interpretations it is desirable to have as precision measure a numerical factor which is a function of both the dispersion and the frequency of occurrence. A precision measure which satisfies that condition is the probable error (r). The probable error is defined as that error which is just as likely to be exceeded as not to be exceeded. To describe the probable error, exactly one-half of the positive errors will fall between 0 and $+r$, and similarly one-half of the negative errors will fall between $-r$ and 0. Further, exactly one-half of all the errors will fall within the range $-r$ to $+r$. Stated still another way, one-half of the errors, disregarding signs, will be equal to or less than r . If we can find the value of t such that exactly one-half of the area under the error curve lies between $-t$ and $+t$, we can locate the value of r for any set of measurements for which σ is known. In Exercise 13 of Chapter V, the value of t corresponding to $\Phi(t) = 0.50000$ was obtained as an illustration of inverse interpolation. The solution gave

$$t_{1/2} = 0.4769$$

Therefore

$$\begin{aligned}\frac{r}{\sqrt{2}\sigma} &= 0.4769 \\ r &= 0.4769 \sqrt{2}\sigma \\ r &= 0.6745\sigma\end{aligned}\tag{27}$$

In equation (27) it has been assumed that σ represents the root-mean-square average of the true errors. We are able to calculate only the standard deviation from the arithmetic mean. Since the principle of least squares shows that the sum of the squares of the deviations from a is a minimum (i.e., less than for any other value) and since the "true value" is only approached by the arithmetic mean, statisticians have suggested that σ for the true errors will be greater than σ for d_a . It has been recommended, therefore, in order to make a better estimate of the probable error, that r should be computed from Bessel's formula

$$r = 0.6745 \sqrt{\frac{\Sigma d_a^2}{n-1}}\tag{28}$$

thereby giving the radical a slightly larger value. When the number of measurements is large, the subtraction of 1 from n will make little difference. A safe rule will consist in observing whether the factor $(n-1)$ instead of n will affect the significant figures of the computed result. If σ is computed by taking the deviations from a , squaring, and summing, it is a simple matter to divide by $(n-1)$ instead of n before taking the square root. If a and σ have been computed by the short-cut method of Chapter VIII, the application of the correction would give

$$r = 0.6745 \sigma \sqrt{\frac{n}{n-1}}\tag{29}$$

An approximation formula for the computation of the probable error, known as Peter's formula, is

$$r = 0.8454\delta\tag{30}$$

Substituting deviations for errors may be supposed to require the same type of correction already seen in the Bessel formula, so that Peter's formula becomes

$$r = 0.8454 \frac{\Sigma |d_a|}{n-1}\tag{31}$$

Peter's formula gives an erroneous result if the measurements are very discordant; hence it is not to be highly recommended.

Verification of the "Law" of Errors. Much has been said about distribution curves for measurements and for errors. It has also been stated that the curve of error is not only a probability curve but a frequency-coefficient curve as well. Before we leave the subject perhaps it will be reassuring to see in just what magic way it is possible for the expression

$$\Phi(t) = \frac{1}{\sqrt{\pi}} \int_{-t}^t e^{-t^2} dt$$

to contain so very much information about a series of measurements. The area under the curve between $-t_1$ and t_1 gives the probability that an error will fall between the limits $-x_1$ and x_1 . The probability that an error will fall between $-x_2$ and x_2 is given by the area under the curve between $-t_2$ and t_2 . If x_2 is greater than x_1 , the probability that an error will lie between $-x_2$ and $-x_1$ or between x_1 and x_2 will be given by the difference between the two areas, or

$$P = \Phi(t_2) - \Phi(t_1) \quad (32)$$

Since the probability is the same as the frequency coefficient, the probability multiplied by the total number of errors gives the frequency of the particular error x . Or

$$n\Phi(t_1) = f_1 \quad (33)$$

and

$$n[\Phi(t_2) - \Phi(t_1)] = f_2 \quad (34)$$

In general

$$n[\Phi(t_i) - \Phi(t_{i-1})] = f_i \quad (35)$$

A group of 76 observations made by Rowland,⁵ between 11° and 20° C, of the mechanical equivalent of heat expressed in kilogram-meters per 10° C, gave an arithmetic mean of $a = 427.2$ and a standard deviation of $\sigma = 0.565$. The distribution of the deviations from a are given in class intervals of 0.2 disregarding plus and minus signs. Assuming that the deviations may be treated as

⁵ Rowland: "On the Mechanical Equivalent of Heat," *Proc. Am. Acad. Arts Sci.*, 15, 75-200 (1880).

errors, the column of calculated frequencies was obtained, using equations (33), (34), and (35).

$\pm d_a$	f (obs.)	$\frac{x}{\sqrt{2}\sigma}$	PROBABILITY	f (calc.)
0.0 to 0.2	22	0.25	0.276	21
.2 " .4	20	0.50	.245	19
.4 " .6	13	0.75	.190	14
.6 " .8	11	1.00	.132	10
.8 " 1.0	6	1.25	.080	6
1.0 " 1.2	1	1.50	.043	3
1.2 " 1.4	2	1.75	.021	2
over 1.4	1	∞	.013	1

The agreement between the observed and calculated frequencies is very good considered in the light of our having treated 76 deviations as if they were an infinite number of errors.

To facilitate an understanding of the above table, let us perform in detail a few of the calculations. The probability that a positive or negative error will lie between 0 and 0.2 is given by the area (from Table II of the Appendix, retaining only three figures) between $t = 0$ and

$$t = \frac{x}{\sqrt{2}\sigma} = \frac{0.2}{\sqrt{2} \cdot 0.565} = \frac{0.2}{0.8} = 0.25$$

From the table,

$$\Phi(0.25) = 0.276$$

Then

$$f = n\Phi(0.25) = 76 \times 0.276 = 21 \text{ (approx.)}$$

The probability of an error falling between 0.2 and 0.4 is given by the area between $t = 0.25$ and

$$t = \frac{0.4}{0.8} = 0.50$$

Then

$$f = n[\Phi(0.50) - \Phi(0.25)] = 76(0.521 - 0.276)$$

$$f = 19 \text{ (approx.)}$$

The remainder of the table is calculated in the same fashion.

By assuming that half of the errors in each class interval were positive and half were negative, it would be possible from the values of n , a , and σ to reconstruct a frequency table of the measurements themselves that would be approximately correct.

PROBLEM SET XIII

1. The following frequency table shows the so-called J-type of distribution. Make graphs to show that the distribution curve of the measurements is of the same form as the distribution curve of the deviations from the arithmetic mean.

x	f	x	f
126	2	129	18
127	5	130	28
128	12		

2. What is the probability of obtaining the king of hearts on a single draw of one card from a deck of 52 playing cards?

3. What is the probability of drawing an ace from the deck of playing cards?

4. What is the probability of not drawing an ace from the deck of playing cards?

5. A problem is assigned to a class of four students who work independently. If each student is just as likely to solve the problem as not to solve it, what is the probability (a) that all four will solve the problem and (b) that the problem will be solved? *Hint:* The problem will be solved unless all four fail.

6. What is the probability of obtaining the sum of (a) 8, (b) 11, and (c) 2, when two dice are thrown?

7. How many three-element products can be made from a third-order determinant?

8. When 10 coins are tossed at once, what is the probability of obtaining (a) exactly 4 heads and (b) at least 4 heads?

9. What is the probability of obtaining exactly 4 heads in both of 2 successive trials in which 10 coins are tossed at once?

10. By means of equation (26), compute to five significant figures the value of $\Phi(t)$ at $t = 1/10$.

11. Making use of the answers to the indicated problems in Set XII, compute the precision (r) of each of the following sets of data: (a) atomic weight of carbon, Baxter and Hale, Problem 1; (b) atomic weight of carbon, Scott and Hurley, Problem 4; (c) atomic weight of chlorine, Hönigschmid and Chan, Problem 6; and (d) ratio of AgI to AgCl, Hönigschmid and Striebel, Problem 7.

CHAPTER X

STATISTICAL INTERPRETATION OF MEASUREMENTS

& þat is a poynt o sorquydryȝe,
þat vche god mon may euel byseme,
To leue no tale be true to tryȝe,
Bot þat hys one skyl may deme.

—*Selection from "The Pearl."*¹

As the result of the extensive treatment of the Gaussian theory of normal distribution of errors in Chapter IX the following important conclusions may be drawn:

1. The most probable representation of central tendency in a set of discordant measurements of a quantity which were performed with equal exactness is the arithmetic mean (a).
2. The best representation of the degree of consistency (precision) in a set of discordant measurements of a quantity which were performed with equal exactness is the probable error (r).

Recalling the meanings of a and r , we can draw further conclusions. Suppose that, for a given set of experimental conditions, the repeated measurement of a designated physical quantity results in stated values of a and r . If more measurements are made of the same quantity under the same conditions, 50 per cent of the further measured values will fall within the range given by

$$a \pm r \quad (1)$$

¹ An English poem of the fourteenth century; authorship unknown. The quotation may be rendered:

And this is a matter of conceit
That ill becomes a good man,—
To believe no tale to be true
Unless his own reason can judge of it.

In other words, the r we have been discussing is the probable error of a single measurement. That is,

$$r \text{ (single observation)} = 0.6745 \sqrt{\frac{\sum d_a^2}{n-1}} \quad (2)$$

We should like to be able to draw some more conclusions in order to complete the picture. We have said that a is the most probable value. How probable? We have said a single measurement has a 50 per cent chance of falling within $a \pm r$. Within what range will 75, 90, and 99 per cent, respectively, of the observations fall? The answers to these questions will have to be deferred until we consider a little more of the indispensable mathematical background.

PROBABLE ERROR OF A COMPUTED RESULT

Suppose the ultimate purpose in measuring the value of a quantity x to be the calculation of y . If

$$y = f(x)$$

an uncertainty in the value of x due to error will produce an uncertainty in the computed value of y . Let us assume that the true value of x is known so that the error may be denoted by Δx . Our mathematics is such that we can reason only in terms of errors and not deviations. Graphically, the values of $f(x)$ may be represented by some curve such as is depicted in Fig. 13. Thus if Δx is the error in x , Δy will be the error of the computed result. The magnitude of the error has been deliberately magnified in the diagram. In practically all cases the errors will be much smaller in comparison with x and y . If the errors are small, the ratio of the errors will approach the value of the slope of the curve at x . Or

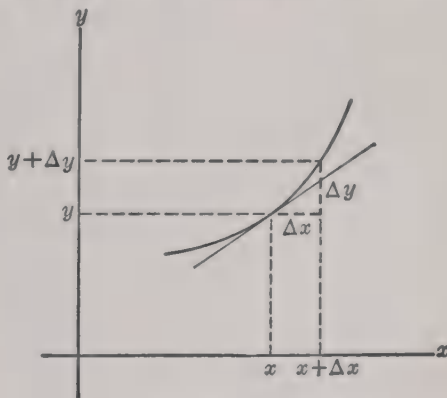


FIG. 13.

$$\frac{\Delta y}{\Delta x} \text{ approaches } \left. \frac{dy}{dx} \right|_{x=i}$$

As a good approximation we can state

$$dy_i = \left(\frac{dy}{dx} \right) dx_i \quad (3)$$

The differentials dy_i and dx_i stand for the errors in y and x , respectively.

In order to obtain a result that will be of more than limited applicability we shall need to extend the argument to cases in which y may be a function of more than one variable, or

$$y = f(x, z, w \dots)$$

Invoking the fundamental principles of partial differentiation, a statement analogous to (3) will be

$$dy_i = \left(\frac{\partial y}{\partial x} \right) dx_i + \left(\frac{\partial y}{\partial z} \right) dz_i + \left(\frac{\partial y}{\partial w} \right) dw_i + \dots \quad (4)$$

Squaring both sides of (4)

$$\begin{aligned} (dy_i)^2 &= \left(\frac{\partial y}{\partial x} \right)^2 (dx_i)^2 + \left(\frac{\partial y}{\partial x} \right) \left(\frac{\partial y}{\partial z} \right) dx_i dz_i + \left(\frac{\partial y}{\partial x} \right) \left(\frac{\partial y}{\partial w} \right) dx_i dw_i \\ &+ \left(\frac{\partial y}{\partial x} \right) \left(\frac{\partial y}{\partial z} \right) dx_i dz_i + \left(\frac{\partial y}{\partial z} \right)^2 (dz_i)^2 + \left(\frac{\partial y}{\partial z} \right) \left(\frac{\partial y}{\partial w} \right) dz_i dw_i \\ &+ \left(\frac{\partial y}{\partial x} \right) \left(\frac{\partial y}{\partial w} \right) dx_i dw_i + \left(\frac{\partial y}{\partial z} \right) \left(\frac{\partial y}{\partial w} \right) dz_i dw_i + \left(\frac{\partial y}{\partial w} \right)^2 (dw_i)^2 \\ &+ \dots \end{aligned} \quad (5)$$

According to the Gaussian theory an error is equally likely to be positive or negative. The squared terms will be independent of sign but the cross products will have varying signs. Hence, if corresponding sets of values of y , x , z , and w are taken and, if the number of sets be sufficiently large, it follows that when the equations are summed the cross products will cancel and only the squared terms will remain. (For typographical convenience the cross products are omitted in the summing operation following.)

$$(dy_1)^2 = \left(\frac{\partial y}{\partial x}\right)^2 (dx_1)^2 + \left(\frac{\partial y}{\partial z}\right)^2 (dz_1)^2 + \left(\frac{\partial y}{\partial w}\right)^2 (dw_1)^2 + \dots$$

$$(dy_2)^2 = \left(\frac{\partial y}{\partial x}\right)^2 (dx_2)^2 + \left(\frac{\partial y}{\partial z}\right)^2 (dz_2)^2 + \left(\frac{\partial y}{\partial w}\right)^2 (dw_2)^2 + \dots$$

$$\begin{array}{cccc} \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot \end{array}$$

$$(dy_n)^2 = \left(\frac{\partial y}{\partial x}\right)^2 (dx_n)^2 + \left(\frac{\partial y}{\partial z}\right)^2 (dz_n)^2 + \left(\frac{\partial y}{\partial w}\right)^2 (dw_n)^2 + \dots$$

$$\Sigma(dy)^2 = \left(\frac{\partial y}{\partial x}\right)^2 \Sigma(dx)^2 + \left(\frac{\partial y}{\partial z}\right)^2 \Sigma(dz)^2 + \left(\frac{\partial y}{\partial w}\right)^2 \Sigma(dw)^2 + \dots \quad (6)$$

Multiplying through by $\frac{(0.6745)^2}{n}$, the result will be

$$R^2 = \left(\frac{\partial y}{\partial x}\right)^2 r_x^2 + \left(\frac{\partial y}{\partial z}\right)^2 r_z^2 + \left(\frac{\partial y}{\partial w}\right)^2 r_w^2 + \dots \quad (7)$$

where R is the probable error in y and the r 's are the probable errors of the indicated variables. Equation (7) is subject to the limitations of the general partial differential formula on which it is based, namely, the variables $x, z, w \dots$ must be *independent*.

For the special case of $y = f(x)$ the general equation reduces to

$$R = \left(\frac{dy}{dx}\right) r_x \quad (8)$$

For example, the circumference of a circle is given by $C = \pi D$. If the probable error of the measurement of the diameter is r , the probable error of the circumference will be

$$R = \frac{dC}{dD} r_D = \pi r_D$$

Probable Error of a Sum or Difference. If

$$y = x \pm z \quad (9)$$

$$R^2 = \left(\frac{\partial y}{\partial x}\right)^2 r_x^2 + \left(\frac{\partial y}{\partial z}\right)^2 r_z^2$$

$$R^2 = (1)^2 r_x^2 + (\pm 1)^2 r_z^2$$

$$R = \sqrt{r_x^2 + r_z^2} \quad (10)$$

Probable Error of a Product. If y is a function of several variables of the form

$$y = kxz \quad (11)$$

Then

$$R^2 = \left(\frac{\partial y}{\partial x}\right)^2 r_x^2 + \left(\frac{\partial y}{\partial z}\right)^2 r_z^2$$

$$R^2 = (kz)^2 r_x^2 + (kx)^2 r_z^2$$

Or

$$R = k \sqrt{z^2 r_x^2 + x^2 r_z^2} \quad (12)$$

Probable Error of a Quotient. When the relationship between y and the variables upon which it depends is of the form

$$y = k \frac{x}{z} \quad (13)$$

the probable error in y will be given, according to (7), by

$$R^2 = \left(\frac{k}{z}\right)^2 r_x^2 + \left(-\frac{kx}{z^2}\right)^2 r_z^2$$

or

$$R = \frac{k}{z} \sqrt{r_x^2 + \frac{x^2 r_z^2}{z^2}} \quad (14)$$

Probable Error of Powers or Roots. When the value of y depends upon a power of x other than first, a generalized equation would be

$$y = b + kx^n \quad (15)$$

The probable error in y will be given, according to (8), by

$$R = nkx^{n-1} r_x \quad (16)$$

The result, equation (16), is equally applicable for positive and negative values of n whether fractional or integral. A negative value of n would give a negative sign to the calculated probable error. This means that a positive error in x would produce a negative error in y . Since the sign of r_x is unknown, the negative sign so produced should be neglected.

Probable Error of Logs. Functions in which the value of y is related to the value of $\log x$ are of frequent occurrence in chemical computations. In general form, if

$$y = k + n \log x \quad (17)$$

for ready treatment by formula (8), we need to transform to natural logs, or, according to page 20,

$$y = k + 0.4343 n \ln x$$

Then

$$R = \frac{0.4343n}{x} r_x \quad (18)$$

Probable Error of Exponential Functions. Many exponential equations cannot be reduced to a linear log form. The differentiation must then be performed on the exponential form. If

$$y = b + ka^x \quad (19)$$

the probable error in y , according to (8), is

$$R = k (\ln a) a^x r_x$$

which may be expressed in terms of common logs as

$$R = \frac{k}{0.4343} (\log a) a^x r_x \quad (20)$$

Probable Error in Complex Functions. The foregoing illustrations should cover the majority of functions that are likely to be encountered in the computation of results from chemical measurements. If in some instances complex differentiation is required, then complex differentiation must be performed. A student of chemistry who is several years removed from his courses in differential calculus need feel no shame over consulting a reference table of standard differential forms.²

As an example of more complex differentiation we might consider the method of computing atomic weights. The atomic weight

² Excellent treatments of differentiation, with especial application to chemical mathematics, are to be found in Mellor: *Higher Mathematics for Students of Physics and Chemistry*, Longmans, Green, New York, 1919; in Daniels: *Mathematical Preparation for Physical Chemistry*, McGraw-Hill, New York, 1928; and in Partington: *Higher Mathematics for Chemical Students*, Methuen and Co., Ltd., London, 1911.

of iodine has been determined by obtaining the ratio between the weights of equivalent amounts of silver iodide and silver chloride. With known values of the atomic weights for silver and chlorine, the atomic weight of iodine can be computed. We can state:

$$\frac{\text{Weight of silver iodide}}{\text{Weight of silver chloride}} = \frac{\text{Ag} + \text{I}}{\text{Ag} + \text{Cl}}$$

where Ag, I, and Cl denote the respective atomic weights. If we let V represent the ratio of weights, then

$$V = \frac{\text{Ag} + \text{I}}{\text{Ag} + \text{Cl}}$$

Or

$$\text{I} = V\text{Ag} + V\text{Cl} - \text{Ag} = \text{Ag}(V - 1) + V\text{Cl}$$

If the probable errors in the determinations of V , Cl, and Ag are known, we can summarize this information as follows:

$$V \pm r_V$$

$$\text{Cl} \pm r_{\text{Cl}}$$

$$\text{Ag} \pm r_{\text{Ag}}$$

The value of the probable error in the computed atomic weight of iodine would be

$$r_I^2 = \left(\frac{\partial I}{\partial V}\right)^2 r_V^2 + \left(\frac{\partial I}{\partial \text{Ag}}\right)^2 r_{\text{Ag}}^2 + \left(\frac{\partial I}{\partial \text{Cl}}\right)^2 r_{\text{Cl}}^2$$

Or

$$r_I^2 = (\text{Ag} + \text{Cl})^2 r_V^2 + (V - 1)^2 r_{\text{Ag}}^2 + V^2 r_{\text{Cl}}^2$$

Therefore

$$r_I = \sqrt{(\text{Ag} + \text{Cl})^2 r_V^2 + (V - 1)^2 r_{\text{Ag}}^2 + V^2 r_{\text{Cl}}^2}$$

Obtaining Congruity in Measurement. When the value of a property is to be computed from observations of other properties the question arises: "How exactly should each observation be performed to produce a desired degree of exactness in the result?" Obviously, it is useless to perform one observation in a laborious manner if the refinement will not give a better result than a less laborious performance.

As an example, we can consider the cryoscopic determination

of molecular weights by the Beckmann method. For dilute solutions, the molecular weight is given by

$$M = K \frac{w}{DW}$$

where M is the molecular weight of the solute, w is the weight in grams of the solute, W is the weight in grams of the solvent, and D is the lowering of the freezing point of the solvent produced by w grams of solute in W grams of solvent. K is a "theoretical" constant, that is, it is not an empirical constant. The value of K is determined from the molecular weight, the freezing point, and the molar heat of fusion of the pure solvent and from the gas constant, R . All the values have been determined with very high precision, so that for our present purpose we can consider the value of K as free from error.

The probable error of the computed molecular weight would be

$$r_M^2 = \left(\frac{\partial M}{\partial w}\right)^2 r_w^2 + \left(\frac{\partial M}{\partial D}\right)^2 r_D^2 + \left(\frac{\partial M}{\partial W}\right)^2 r_W^2$$

Differentiating, we obtain

$$r_M^2 = \frac{K^2}{D^2 W^2} r_w^2 + \frac{K^2 w^2}{D^4 W^2} r_D^2 + \frac{K^2 w^2}{D^2 W^4} r_W^2$$

If the first term of the right-hand side is multiplied by w^2/w^2 , we can factor the expression to get

$$r_M^2 = \frac{K^2 w^2}{D^2 W^2} \left(\frac{r_w^2}{w^2} + \frac{r_D^2}{D^2} + \frac{r_W^2}{W^2} \right)$$

Clearly, the common factor is equal to M^2 , hence

$$\frac{r_M^2}{M^2} = \frac{r_w^2}{w^2} + \frac{r_D^2}{D^2} + \frac{r_W^2}{W^2}$$

Multiplying through by $(100)^2$

$$E_M^2 = E_w^2 + E_D^2 + E_W^2$$

in which we have let E represent the precision expressed as percentage.

In a typical experiment, the order of magnitude of the observed quantities are: 0.5 g of solute dissolved in 25 g of solvent to

produce a lowering of 0.5° C. Suppose that it is desired to determine a molecular weight with a precision of ± 0.1 per cent. The two weighings can be performed easily with a precision of ± 0.0002 g. In terms of percentage this means 0.04 per cent in the weight of the solute and 0.0008 per cent in the weight of the solvent. The precision with which the temperature can be measured will be the limiting factor. Or

$$E_D^2 = E_M^2 - E_w^2 - E_W^2$$

$$E_D = \sqrt{(0.1)^2 - (0.04)^2 - (0.0008)^2}$$

$$E_D = \sqrt{0.01 - 0.0016 - 0.00000064}$$

$$E_D = \sqrt{0.008} = 0.09\%$$

The temperature change would have to be measured with a precision of ± 0.09 per cent for the determination of the molecular weight to be ± 0.1 per cent. The Beckmann differential thermometer can measure temperature differences with a precision of only $\pm 0.002^{\circ}$ C, which for a one-half degree lowering would be ± 0.4 per cent. It would be necessary to use either a thermocouple or thermopile to attain the higher precision of ± 0.09 per cent.

Let us now answer the converse question: "How precise a result can be obtained by a given method without modification?" Consider again the cryoscopic measurement of molecular weight with the Beckmann thermometer as the means of measuring temperature. The precision of the computed result will be

$$E_M = \sqrt{(0.04)^2 + (0.4)^2 + (0.0008)^2}$$

$$E_M = \sqrt{0.0016 + 0.16 + 0.00000064}$$

The two smaller values are negligible in comparison with the value 0.16, hence

$$E_M = \pm 0.4\%$$

Further, we see that, if the uncertainty in the weighing of the sample of solute is negligible, then the weighing of the sample of solvent can be performed much less exactly. If the precision in the weighing of the solvent were 0.04 per cent, the precision of the computed result would be unaffected. Therefore the weighing of

the solvent need be carried no further than to the nearest hundredth of a gram, that is,

$$25.00 \pm 0.01$$

PROBABLE ERROR OF THE ARITHMETIC MEAN

The arithmetic mean is computed from a series of measured values by use of the formula:

$$a = \frac{x_1 + x_2 + x_3 + \cdots + x_n}{n} \quad (21)$$

Therefore a is a function of each of the x 's. According to equation (7) the probable error of the mean will be given by

$$R^2 = \left(\frac{\partial a}{\partial x_1}\right)^2 r_1^2 + \left(\frac{\partial a}{\partial x_2}\right)^2 r_2^2 + \cdots + \left(\frac{\partial a}{\partial x_n}\right)^2 r_n^2$$

or

$$R^2 = \frac{1}{n^2} r_1^2 + \frac{1}{n^2} r_2^2 + \cdots + \frac{1}{n^2} r_n^2$$

The probable error of a single observation is r , and each single measurement has the same probable error. Therefore,

$$r = r_1 = r_2 = r_3 = \cdots = r_n$$

Then

$$R^2 = \frac{nr^2}{n^2}$$

or

$$R = \frac{r}{\sqrt{n}} \quad (22)$$

We shall designate the probable error of a by r_a in the future. From Bessel's formula (p. 170)

$$r_a = 0.6745 \sqrt{\frac{\sum d_a^2}{n(n-1)}} \quad (23)$$

The probable error of the mean has the following significance: If the experiments by means of which a was found were repeated with the same care and exactness, 50 per cent of the further determinations of a would fall within the range given by

$$a \pm r_a$$

In other words, r_a gives the precision of the arithmetic mean. If the experimental method by which a was determined were known to be free from all sources of corrigible error, then we could state: The "true value" of the given physical quantity has a 50 per cent chance of falling within the range given by

$$a \pm r_a$$

The **if** prefacing this last statement should by merit of relative importance be set in type the size of a page because the whole question of measurement hinges upon it. How is one to know whether all sources of corrigible error have been obviated? One possible answer is to be found in the repeated measurement of a quantity by different observers, preferably using different techniques or experimental methods. Comparisons of several separate determinations should indicate the presence of constant errors. We must now consider how large the differences between separate determinations must be to allow the identification of a constant error which cannot be explained away as a random error. The Gaussian theory, if accepted literally, admits the possibility of a random error infinite in size. In order to be able to identify constant errors we must fix a reasonable limit which random errors are very unlikely to exceed in size. These considerations lead us directly to the concept of a statistically reliable difference.

Statistically Reliable Difference. It is manifestly absurd to credit a physical reality to a random error which has infinite size. In other words, it is unreasonable to postulate a set of experimental conditions under which the imponderable variables of the system happen to combine their effects in such a way as to produce a resultant error which has an infinite magnitude. From the mathematical properties of the curve of error (p. 163) we recall that small errors have high probability and, conversely, large errors have low probability. The problem is to set an arbitrary limit to the smallness of probability beyond which we can say that the probability is negligible. The limit for smallness has been chosen as the probability of an error which is equal to $4r$, or four times the size of the probable error. Let us see what numerical value this vanishingly small probability has. Since

$$r = 0.6745\sigma$$

then

$$\sigma = \frac{r}{0.6745}$$

From page 164, we recall that

$$t = \frac{x}{\sqrt{2} \sigma}$$

When $x = 4r$,

$$t = \frac{4r}{\sqrt{2} \frac{r}{0.6745}} = 1.908$$

From Table II (Appendix), the value of $\Phi(t)$ corresponding to 1.908 is

$$\Phi(1.908) = 0.99302$$

This means that a random error has a 99.3 per cent chance of being equal to or smaller than $4r$. It also means that a random error has only a 0.7 per cent chance of exceeding $4r$. A probability of 0.007 represents 1 chance out of 143. In terms of practical measurements, where the number of trials rarely exceeds 20, this is a vanishingly small probability. If a deviation is observed which is larger than four times the probable error, it has an exceedingly small chance of being a random error. It follows, therefore, that a deviation greater than $4r$ is practically certain *not* to be a random error. A deviation which is not a random error must represent a difference with a real significance. A significant difference might be the result of:

1. The measurement of an entity or quantity which is not the same as that supposedly being measured.
2. The measurement of a quantity when a ponderable variation of the conditions has not been detected.
3. The measurement of a property when a deliberate change of experimental conditions has introduced a new variable, that is, a previously imponderable variable has become ponderable.
4. A mistake.

Whatever the cause, we can state the following conclusion: A deviation which is greater than four times the probable error can be shown by statistical reasoning to represent a significant difference. In brief, a deviation greater than $4r$ is a statistically reliable difference.

The choice of $4r$ as the limit which a random error cannot exceed is, of course, highly arbitrary. Hence this limit cannot

be accepted as an inflexible dictum. Individual cases will require separate consideration.

Interpretation of Statistically Reliable Difference. Let us suppose that we have two values of a which are the means of two sets of observations of a physical quantity. Assume also that the probable errors of the two means have been obtained so that we have the data

$$a_1 \pm r_1$$

$$a_2 \pm r_2$$

where for momentary convenience the subscript a 's have been omitted from the r 's. Denote the difference between the means by v ; then

$$v = a_1 - a_2$$

and by equation (10) the probable error of the difference is

$$R = \sqrt{r_1^2 + r_2^2}$$

If the difference v is equal to or just greater than $4R$, there is only 1 chance in 143 that v can be attributed to random error. If v is greater than R by a factor of 5, 6, 7, etc., the probability becomes much less. (See Problem 11 at the end of this chapter.) If the difference between the two means is not attributable to the random operation of the imponderable variables of the system, at least two interpretations are possible:

1. The two sets of measurements represent observations of two different things.
2. One or both sets of measurements are subject to constant error.

Example. Lord Rayleigh's measurements of the density of nitrogen are given in *Proc. Roy. Soc. (London)*, 55, 340-4 (1894). Nitrogen was prepared from several different sources by various methods. It was found that the observed values for the mass of a fixed volume of gas fell into two distinct groups. The values for nitrogen from the atmosphere agreed well with one another, and the values for nitrogen from nitrogenous compounds were consistent among themselves. The values for the mass of nitrogen which filled a certain flask under the specified conditions of temperature and pressure are as follows, with the method of preparation indicated.

CHEMICAL NITROGEN:

2.30143	}	NO on hot iron.
2.29890		
2.29816		
2.30182		
2.29869	}	N ₂ O on hot iron.
2.29940		
2.29849	}	NH ₄ NO ₂ on hot iron.
2.29889		
2.30074	}	N ₂ O on electrified, hot iron.
2.30054		

ATMOSPHERIC NITROGEN:

2.31017	}	Air on hot iron.
2.30986		
2.31010		
2.31001		
2.31024	}	Air with ferrous hydroxide.
2.31010		
2.31028		
2.31163	}	Air on electrified, hot iron.
2.30956		

For **chemical nitrogen**, the mean and its probable error are

$$2.29971 \pm 0.00028$$

For **atmospheric nitrogen** the corresponding values are

$$2.31022 \pm 0.00013$$

The difference between the two means is

$$v = 2.31022 - 2.29971 = 0.01051$$

The probable error of the difference, v , is

$$R_v = \sqrt{(2.8 \times 10^{-4})^2 + (1.3 \times 10^{-4})^2}$$

$$R_v = 0.00031$$

Hence

$$\frac{v}{R_v} = \frac{0.0105}{0.00031} = 34.$$

From this it is easy to see that the difference is far more than four times its probable error. Therefore there is a statistically reliable difference between the two sets of observations. The wide variety of methods employed in preparing nitrogen should have eliminated the possibility of

there being a constant error due to manipulation. Lord Rayleigh made no attempt to suppress these discordant data. On the contrary, he emphasized the discrepancy. It was just this emphasis which led to the discovery in the atmosphere of the inert gases (the so-called noble gases, He, Ne, A, etc.) by Sir William Ramsay.³

In the light of our present chemical knowledge, we should be inclined to speak of the presence of an impurity as the source of a constant error. Hence we see that we cannot draw a sharp distinction between "observations of different entities" and "observations subject to constant error."

Often, the object of research may be to establish whether the deliberate introduction of a new variable produces a statistically reliable difference in the value of a property. For instance, a chemical reaction rate might be measured both in the absence of and in the presence of an alternating electromagnetic field. The test for a statistically reliable difference when applied to the two results could tell whether the intensity of electromagnetic radiation of given frequency is a ponderable variable for the particular reaction.

When a statistically reliable difference has been established and an explanation for the difference has not been found either theoretically or experimentally, statistical considerations can throw little further light on the problem. We have no mathematical criterion for choosing the correct mean and rejecting the one in error. We cannot even be sure that one is right and the other wrong. If the probable errors of the two means are widely different we should give more credence to the mean having the smaller probable error. The only reliable solution of the difficulty lies in repetition of the measurement, preferably by several observers and by more than one experimental method.

REJECTION OF SUSPECTED OBSERVATIONS

Often, in a series of measurements, a value occurs which deviates widely from the remainder of the values. If the widely divergent value is retained it will produce a marked effect upon the values of the arithmetic mean and probable error which represent the whole series of measurements. A beginner will be strongly tempted to discard the one "bad" value as being obvi-

³ See Rayleigh and Ramsay: *Phil. Trans.*, **186(A)**, 187-241 (1895).

ously wrong and thereby greatly enhance the appearance of his results. A golfer who neglects to record on his score card the extra stroke required to get out of a sand-trap is considered dishonest and a poor sportsman. An experimenter who unreasoningly discards a discordant value as "obviously wrong" is dishonest and a poor scientist. If the experimenter can prove conclusively to himself that a mistake was made such as: (a) the incorrect numerical entry of a scale length or a sum of weights, etc.; (b) the accidental presence of a visible contamination; (c) the accidental loss of material; or (d) some other equally valid mistake, he is completely justified in discarding the bad value. Failing to establish such proof, he may seek guidance from the concept of a statistically reliable difference.

We have said that, if a further observation of a previously measured quantity deviates from the mean of the results by more than four times the probable error of a single observation, the deviation is very unlikely to be a random error. The limit $4r$ was presumed for an indefinite, small number of measurements. If the causative factors to which random errors are attributed are recalled from previous discussions, the more times a measurement is performed, assuming identical conditions for each observation, the greater will be the probability of the proper combination of circumstances to produce large random errors. Hence the assumption of a fixed limit (i.e., $4r$) appears to be somewhat crude to serve as a criterion in the delicate task of rejecting observations. Various criteria have been evolved, taking into account the probability aspect of the number of observations. The criterion which is generally considered to be most worthy of notice is the one due to Pierce and Chauvenet. We shall not discuss the details;⁴ we need consider only the results. The criterion and the procedure of application are as follows:

1. Compute the mean and probable error of a single observation, retaining all suspected observations in the computation.
2. Determine the ratio of the suspiciously large deviation to the probable error of a single observation.
3. From the following table secure the limiting value of the ratio for the corresponding number of observations (n).

⁴ For a detailed analysis of the Pierce-Chauvenet criterion, see Chauvenet: *Spherical and Practical Astronomy*, Vol. II, p. 558, Lippincott, Philadelphia, 1868.

4. If the observed value of d/r is greater than the tabular value of d/r , the observation may be rejected.

PIERCE-CHAUVENET CRITERION

n	$\frac{d}{r}$
5	2.5
10	2.9
15	3.2
20	3.3
50	3.8
100	4.2

Example. The iron content of an aluminum casting alloy was determined by a group of supposedly equally reliable analysts employing a variety of volumetric, gravimetric, and colorimetric methods. The results, expressed as percentage of iron, tabulated together were as follows. Should the value 1.83 be rejected?

% Fe	d_a	d_a^2
1.52	-0.044	0.001936
1.46	- .104	10816
1.61	+ .046	2116
1.54	- .024	576
1.55	- .014	196
1.49	- .074	5476
1.68	+ .116	13456
1.46	- .104	10816
1.83	+ .266	70756
1.50	- .064	4096

$$a = 1.564$$

$$0.120240$$

$$r = 0.6745 \sqrt{\frac{\Sigma d_a^2}{n-1}} = 0.6745 \sqrt{\frac{0.1202}{9}} = 0.078$$

The ratio of the suspected deviation to the probable error of a single observation is

$$\frac{d}{r} = \frac{0.266}{0.078} = 3.4$$

For ten observations the tabulated value is 2.9. Since

$$3.4 > 2.9$$

the observation 1.83 should be rejected.

At best, this criterion can be recommended only as a scheme for the guidance of the inexperienced to prevent thoughtless rejection of values. The soundest recommendation, though not always the most practical, is that enough measurements be performed to overcome the effect of the one highly discordant value. The experimenter can doubtless find some consolation in the assumption that, if the large deviation is a random error, the possibility exists that a sufficient number of trials may produce an equally large deviation of opposite sign.

STATISTICALLY WEIGHTED AVERAGE

Very often several independent measurements of a physical quantity will have been made either by different observers or by one observer employing different methods. The values of a and r_a having been obtained for each set of measurements, the question arises as to how these independent values of a shall be combined to yield the best representative average. From the results of the last chapter we recall that

$$P = \frac{1}{\sqrt{2\pi} \sigma} e^{-\frac{x^2}{2\sigma^2}} \quad (24)$$

gives the probability of an error x when its measure of dispersion is σ . If the values of a are assumed to follow the normal law of distribution, then

$$P_1 = \frac{1}{\sqrt{2\pi} \sigma_1} e^{-\frac{(a_1 - z)^2}{2\sigma_1^2}}$$

gives the probability of the given value a_1 , if z is assumed to be the most probable representative value and σ_1 is the dispersion measure of a_1 . Likewise the probability of a_n will be

$$P_n = \frac{1}{\sqrt{2\pi} \sigma_n} e^{-\frac{(a_n - z)^2}{2\sigma_n^2}}$$

The values of the a 's were independently determined; hence the probability of obtaining the set of n values of a will be the product of the separate probabilities, or

$$P = \frac{1}{\sqrt{2\pi} \sigma_1} e^{-\frac{(a_1 - z)^2}{2\sigma_1^2}} \cdot \frac{1}{\sqrt{2\pi} \sigma_2} e^{-\frac{(a_2 - z)^2}{2\sigma_2^2}} \cdots \frac{1}{\sqrt{2\pi} \sigma_n} e^{-\frac{(a_n - z)^2}{2\sigma_n^2}}$$

For P to be maximum, the sum of the exponents must be a minimum, or

$$Q = \frac{(a_1 - z)^2}{2\sigma_1^2} + \frac{(a_2 - z)^2}{2\sigma_2^2} + \dots + \frac{(a_n - z)^2}{2\sigma_n^2}$$

must be a minimum. The necessary condition for Q to be minimum is

$$\frac{dQ}{dz} = 0$$

Then

$$-\frac{2(a_1 - z)}{2\sigma_1^2} - \frac{2(a_2 - z)}{2\sigma_2^2} - \dots - \frac{2(a_n - z)}{2\sigma_n^2} = 0$$

from which

$$z \left(\frac{1}{\sigma_1^2} + \frac{1}{\sigma_2^2} + \dots + \frac{1}{\sigma_n^2} \right) = \frac{a_1}{\sigma_1^2} + \frac{a_2}{\sigma_2^2} + \dots + \frac{a_n}{\sigma_n^2}$$

dividing both sides by $(0.6745)^2$

$$z \left(\frac{1}{r_1^2} + \frac{1}{r_2^2} + \dots + \frac{1}{r_n^2} \right) = \frac{a_1}{r_1^2} + \frac{a_2}{r_2^2} + \dots + \frac{a_n}{r_n^2}$$

or

$$z = \frac{\frac{a_1}{r_1^2} + \frac{a_2}{r_2^2} + \dots + \frac{a_n}{r_n^2}}{\frac{1}{r_1^2} + \frac{1}{r_2^2} + \dots + \frac{1}{r_n^2}} \quad (25)$$

In the light of equation (5), page 127, it is apparent that the above equation is an expression of a weighted arithmetic mean in which the values of a are weighted according to the inverse squares of their probable errors. The r 's are values of r_a . In more general form

$$z = \frac{\sum \frac{a}{r_a^2}}{\sum \frac{1}{r_a^2}} \quad (26)$$

If the deviations from the weighted mean (z) are obtained

$$d_i = a_i - z$$

then

$$r_z = 0.6745 \sqrt{\frac{\frac{d_1^2}{r_1^2} + \frac{d_2^2}{r_2^2} + \cdots + \frac{d_n^2}{r_n^2}}{\frac{1}{r_1^2} + \frac{1}{r_2^2} + \cdots + \frac{1}{r_n^2}}} \cdot \frac{1}{n-1} \quad (27)$$

gives the probable error of the weighted mean. In general, the probable error of any weighted mean is given by

$$R = 0.6745 \sqrt{\frac{\sum w d^2}{(n-1) \sum w}} \quad (28)$$

We can obtain a different estimate of the probable error of z [equation (25)] by means of the partial differential expression (7). The probable error of z would be

$$R^2 = \left(\frac{\partial z}{\partial a_1} \right)^2 r_1^2 + \left(\frac{\partial z}{\partial a_2} \right)^2 r_2^2 + \cdots + \left(\frac{\partial z}{\partial a_n} \right)^2 r_n^2$$

or

$$R^2 = \left(\frac{\frac{1}{r_1^2}}{\sum \frac{1}{r^2}} \right)^2 r_1^2 + \left(\frac{\frac{1}{r_2^2}}{\sum \frac{1}{r^2}} \right)^2 r_2^2 + \cdots + \left(\frac{\frac{1}{r_n^2}}{\sum \frac{1}{r^2}} \right)^2 r_n^2$$

whence

$$R^2 = \frac{\sum \frac{1}{r_a^2}}{\left(\sum \frac{1}{r_a^2} \right)^2} = \frac{1}{\sum \frac{1}{r_a^2}}$$

or

$$r_z = \frac{1}{\sqrt{\sum \frac{1}{r_a^2}}} \quad (29)$$

Denoting r_z obtained by equation (27) as the "observed" probable error and the r_z computed from equation (29) as the "theoretical" probable error, we can write the expression:

$$\text{Ratio} = \frac{\text{Observed } r_z}{\text{Theoretical } r_z}$$

If the precision measures of the various a 's were correct and there were no constant errors, the above ratio would be exactly unity. In practice a ratio between $\frac{1}{2}$ and 2 may reasonably be expected. If the ratio is much larger than 2, the data are probably subject to constant errors. If constant errors are present, the set of means may be critically examined with the view of possible rejection of a highly discordant value. If a well-justified rejection reduces the ratio to a reasonable value, the weighted mean may be accepted with some measure of confidence. Otherwise, the data ought never to have been averaged together in the first place because statistical treatment, based as it is on probability, has no application to data of measurement where there are constant errors as large as or larger than the random errors.

LIMITED NUMBER OF OBSERVATIONS

The Gaussian theory of errors was derived for an infinite number of observations. We have seen that it yields predictions which are in good agreement with practical cases when the number of observations is reasonably large. We should expect to find that the statistical method will be less rigorously applicable as the number of observations decreases to small values. We shall not be surprised to find, therefore, that a precision measure based upon only four or five deviations is less reliable than one based upon a much larger number of values. H. Jeffreys⁵ has proposed a modification of the usual calculation of probable errors which purports to yield precision measures that are more reliable. Our formulas are

$$r = 0.6745 \sqrt{\frac{\sum d_a^2}{n-1}}$$

and

$$r_a = 0.6745 \sqrt{\frac{\sum d_a^2}{n(n-1)}}$$

The following table gives Jeffreys' values of the constant which should replace the 0.6745 of the above equations when the number of observations is very small.

⁵ Jeffreys: *Proc. Roy. Soc. (London)*, **138A**, 48 (1932).

n	JEFFREYS' CONSTANT
2	1.000
3	0.816
4	0.766
5	0.740
6	0.728
7	0.718
10	0.703
∞	0.674

For values of n greater than 10, the usual constant gives sufficiently reliable estimates.

The probable error determined from two observations is a very vague attribute of those measurements. One positive and one negative deviation which are of necessity identical in magnitude seem a poor basis upon which to predict limits between which 50 per cent of all other observations of that quantity will fall. In the first place, one chief objective in evolving a precision measure is to provide a lucid, shorthand method of stating results of measurements. When there are only two values, the excuse for a precision measure is removed. A statement of the two values and their mean will serve amply to describe the results of the experiment. When there are three or four measurements, a mean value begins to have some significance, and, even though the necessity for a condensed statement is not imperative, there is sufficient certainty in the value of the probable error to make its computation worth while.

NUMERICAL STATEMENTS OF PRECISION MEASURES

Up to this point, we have stated precisions, errors, deviations, and uncertainties almost entirely as *absolute statements of value*. That is, the magnitude of the measure has been expressed along with the representative value of the given set of data. Examples are:

1. Uncertainty of a number 1.064 ± 0.001
2. Probable error of single observation $a \pm r$
3. Probable error of mean $a \pm r_a$
4. Standard deviation $a \pm \sigma$
5. Average deviation $a \pm \delta$

In addition to *absolute* statements of precision, for purposes of comparison, it is convenient to state *relative* values of precision,

etc. Each of the three commonly employed methods of expressing relative precision, etc., will have particular advantages under various conditions of comparison. The three ways of relative expression are:

1. *Fractional.* The given measure of precision, etc., may be stated as a ratio to its representative central value, thus:

$$\pm \frac{r}{a} \quad \text{or} \quad \pm \frac{\sigma}{a}$$

The fractional expression is in the form of a decimal fraction, and hence it is a trifle awkward. The comprehension of a fractional precision or uncertainty by most people requires the mental counting of decimal places. In the end the results are mentally converted to other terms such as percentage. For that reason, the percentage expression is more readily comprehended.

2. *Percentage.* The idea of expressing a ratio as per cent or parts out of a hundred will hardly require elaboration. The percentage precision or uncertainty would be

$$\pm \frac{r_a}{a} 100 \quad \text{or} \quad \pm \frac{0.001}{1.064} 100$$

In a few types of chemical measurement the uncertainty will be of the order of integral values when expressed in parts per 100. In other words, the uncertainty is 1 per cent or greater. When this is so, the percentage method gives a readily comprehensible statement.

3. *Parts per thousand.* Most chemical measurements are performed with a degree of exactness which results in precision measures that would give fractional percentages. For this reason the expression of uncertainties, etc., in parts per 1,000 is usually preferable. In routine methods of quantitative analysis the agreement to be expected is 1 to 3 parts per 1,000.

It will be recalled from our treatment of significant figures in Chapter II that the method of avoiding non-significant zeros in numerical values involved transposition to other, more convenient units. The problem here is very much the same. We are confronted with a choice of the "unit" in which precisions, etc., are to be expressed. For convenience we usually strive for a "properly expressed digit factor" (p. 2). It may often be convenient

to choose parts per ten thousand, or, in some highly refined physical measurements, parts per million will constitute the proper choice.

The statement of "parts per thousand" must not be confused with "parts in the thousandths place." Thus, if a properly expressed numerical value is

$$4.893$$

the uncertainty will be ± 0.001 . The uncertainty is 1 in the thousandths place, but on the other hand it is 1 part per 5,000.

SIGNIFICANT FIGURES

By this time, undoubtedly, the reader will have explained to himself the reasoning back of the statements of the computation rules for significant figures which were laid down so dogmatically in Chapter II—all except perhaps Rules 9 and 10. We had better consider these two rules in reverse order.

Rule 10 states: "In a precision measure, retain only two significant figures." We shall have to examine the question of the precision of the precision measure. That means the probable error of the probable error. This looks at first glance like statistical considerations prolonged *ad absurdum* or perhaps even *ad nauseum*. Nevertheless, completely sound statistical reasoning shows that the probable error of the probable error is approximately

$$\pm \frac{0.48}{\sqrt{n-1}} r$$

where n is the number of measurements upon which the value of r is based. When n is 4, the uncertainty in the value of r is 29 per cent. If there are ten measurements, the uncertainty of r is only 16 per cent. When n is as great as 50, the uncertainty drops to 7 per cent. We perceive that with such large uncertainties the second figure of the precision estimate will be uncertain, hence any figures beyond the second can have no significance unless the number of measurements is very large. When the number of measurements is extremely small, there may be no certain figures in the value of the probable error. Hence, in those cases, only one figure may be considered significant.

Rule 9 states that one extra figure should be carried in the mean of four or more values. There is no statistical basis for this

rule; it is merely a dictum which has become apparent from experience. It is often true that the magnitude of the probable error of the mean is such that the number of significant figures in the mean is one greater than the number in the measured values. The reader need only glance back over some of the problems and illustrative examples to see that this is true. In essence, Rule 9 means that one extra figure should be carried in the average unless the precision measure dictates otherwise.

SOME APPROXIMATIONS

In 1830, Auguste Comte in his *Philosophie Positive* said: "If mathematical analysis should ever hold a prominent place in chemistry—an aberration which is happily almost impossible—it would occasion a rapid and widespread degeneration of that science."⁶ This viewpoint was transmitted to chemists and held by them for many years. With theoretical chemistry established upon the "mathematical foundation" where the researches of the last four or five decades have placed it, Comte's viewpoint has become almost as archaic as the phlogiston theory. The fact that mathematical analysis is not rigorously applicable in chemistry now is probably due to our ignorance of chemical phenomena rather than to some fault of mathematics. Chemistry has not yet had a Newton, a Darwin, or an Einstein.

This long-retained idea that mathematics is foreign to chemistry resulted in a feeling of contempt, by many chemists, for anything mathematical. The feeling persists even today in the guise of the attitude, taken by some, that chemistry is largely an inexact science and hence almost any approximation that greatly reduces computation is "good enough." It often requires a very wise person to decree what is "good enough" and thereby secure consistency. Furthermore, the prosecution of the "good enough" attitude to the extreme can be very detrimental to the progress of a science and to the understanding of a beginner.

In many textbooks of chemistry, notably those on quantitative analysis, approximations are given for the precision measure but unfortunately the approximate nature is not emphasized. As a consequence the highly approximate considerations are unintelli-

⁶ Cf. Mellor, *op. cit.*, p. 4:

gently applied far beyond the range over which they are good enough.

The approximate precision measure which is recommended is the average deviation instead of the probable error. If the number of measurements is very few, the numerical value of δ is quite near the value of r as calculated from Peter's formula (p. 170). Thus, for three measurements,

$$\delta = \frac{\Sigma |d_a|}{3} = 0.33 \Sigma |d_a|$$

$$r = 0.85 \frac{\Sigma |d_a|}{3 - 1} = 0.43 \Sigma |d_a|$$

Or for five measurements

$$\delta = \frac{\Sigma |d_a|}{5} = 0.20 \Sigma |d_a|$$

$$r = 0.85 \frac{\Sigma |d_a|}{5 - 1} = 0.21 \Sigma |d_a|$$

However, for twenty measurements,

$$\delta = \frac{\Sigma |d_a|}{20} = 0.050 \Sigma |d_a|$$

$$r = 0.85 \frac{\Sigma |d_a|}{20 - 1} = 0.044 \Sigma |d_a|$$

For very few measurements the discrepancies are of the order of magnitude of the uncertainties in the values of the probable errors. Values of r from Peter's formula are, however, only approximations to Bessel's formula (p. 170), so that the above agreement is an approximate agreement with an approximation. For larger numbers of measurements, it is erroneous to call the average deviation the precision measure.

Another approximation formula gives the precision of the arithmetic mean as

$$a \pm \frac{\delta}{\sqrt{n}}$$

This is subject to all the above limitations.

A simplified criterion for rejection of suspected observations is as follows:

1. Compute the mean and the average deviation of all the measurements, omitting the doubtful value.
2. If the deviation of the suspected value from the mean is as great as or greater than 4δ , the measurement may be rejected.

This criterion is an almost direct invocation of the statistically reliable difference concept presuming the approximate equality of r and δ . For very few measurements, δ is smaller than r ; above six measurements δ becomes larger than r . Fortunately, this criterion can be extended to large numbers of measurements with less harmful effects than can the preceding approximations. For large values of n , this criterion is more stringent than Chauvenet's. As a result the beginner, in using this approximation, will be less likely to reject borderline values indiscriminately.

In that these approximations are truly "good enough" for interpretation in routine quantitative analysis, the justification for their limited use is sufficient. The chief objection that can be raised is one of consistency. There is no reason for calling John Jones "John" other than that his name is John.

PROBLEM SET XIV

1. Avogadro's number N is evaluated by this equation

$$N = \frac{Fc}{10e}$$

where F is the Faraday, c is the velocity of light, and e is the electronic charge.
Given:

$$F = 96,489 \pm 7 \text{ coulomb per equiv}$$

$$c = (2.99796 \pm 0.00004) \times 10^{10} \text{ cm per sec}$$

$$e = (4.770 \pm 0.005) \times 10^{-10} \text{ electrostatic unit}$$

Compute the value of N and its probable error.

2. The 1939 international atomic weight of hydrogen⁷ is 1.00805 ± 0.00002 . Using the value of N obtained in Problem 1, compute the mass of an atom of hydrogen and its probable error.

3. From the results of Problems 1 and 4, Set XII, and of Problem 11 (a) and (b), Set XIII, compute the probable error of the mean atomic weight of carbon as determined (a) by Baxter and Hale and (b) by Scott and Hurley.

⁷ Cf. *Rev. Mod. Phys.*, **9**, 370 (1937).

4. Using the answers of Problem 3 (above), determine the statistically weighted mean atomic weight of carbon and the weighted probable error. Why do you suppose the international atomic weight table for 1938 gives $C = 12.010$?

5. Compute the values of r_a for (a) the atomic weight of chlorine [Problem 6, Set XII, and Problem 11 (c), Set XIII] and (b) the ratio of silver iodide to silver chloride [Problem 7, Set XII, and Problem 11 (d), Set XIII].

6. The atomic weight of silver is

$$\text{Ag} = 107.880 \pm 0.001$$

Employing the results of Problem 5 (above) and of the example on page 180, compute the atomic weight of iodine and its probable error.

7. From the data of the example on page 190, compute the values of a and r for the nine observations remaining after the rejection of the value 1.83. Apply the Pierce-Chauvenet criterion to determine whether the value 1.68 should be rejected.

8. In a certain aqueous solution the mean value of the hydrogen-ion concentration was $(2.412 \pm 0.011) \times 10^{-6}$. Compute the pH and its probable error.

9. From the following six independent determinations of Planck's constant, h , compute the statistically weighted mean and the probable error of the weighted mean by two methods. Obtain the ratio of the two probable errors and interpret. (See p. 193.)

SOURCE	$h(\text{erg-sec})$
Rydberg constant	$(6.547 \pm 0.011) \times 10^{-27}$
Ionization potentials	(6.560 ± 0.015) "
X-rays	(6.550 ± 0.009) "
Photoelectric effect	(6.543 ± 0.010) "
Wien's law	(6.548 ± 0.015) "
Stefan-Boltzmann constant	(6.539 ± 0.010) "

10. The mass of the electron as determined from spectroscopic measurements is $(9.035 \pm 0.010) \times 10^{-28}$ g, whereas deflection measurements yield the value $(8.994 \pm 0.014) \times 10^{-28}$ g. Is there a statistically reliable difference between the two values?

11. What is the probability that an error which is (a) $5r$, (b) $6r$, or (c) $10.5r$ can be a random error?

12. The diameter of a sphere was found to be 4.182 ± 0.004 cm. Compute (a) the area of its surface and the probable error, and (b) the volume and its probable error; if the mass of the sphere is 187.4 ± 0.7 g, compute (c) the density and its probable error.

Note. Except where indicated otherwise, the values of the physical constants given in this set of problems are taken from Birge, *Phys. Rev. Suppl.*, **1**, 1 (1929). The student should consult this article and set further exercises for himself from the data given. Also, he should not fail to read an article by Power, treating the errors of microanalysis, in *Ind. Eng. Chem., Anal. Ed.*, **11**, 660-73 (1939).

CHAPTER XI

CURVE FITTING

The assignment of numerals to represent telephones or the articles in a salesman's catalogue is not measurement; nor—and here is a more definite representation of properties—the assignment of numerals to the colours in a dyer's list. The difference between such assignments and true measurement is that they are completely arbitrary. Even true measurement . . . is not completely free from an arbitrary element, but it is based upon something that is not arbitrary at all, namely, scientific laws.

—N. R. CAMPBELL.

In the preceding chapter we have considered the statistical methods of interpreting the results of repeated measurements of a property of a system when all the ponderable variables are fixed or corrected for slight variation. Very often we are interested in the variation of a given property when one of the ponderable variables is successively fixed at various values. The given property becomes the dependent variable, the values of which we measure at various fixed values of the ponderable variable which by definition is the independent variable. Thus, if we measure the volumes occupied by a given mass of a gas at various fixed values of pressure, volume is the dependent and pressure the independent variable. In order to obtain a reasonably complete treatment of the subject of measurement we shall have to consider some of the more frequently used methods of formulating a mathematical relationship between the values of the dependent and independent variables. We shall thereby gain an insight into the idea of "partial laws" discussed on pages 89–91. More times than not, the relationship will be of a purely empirical nature. The formulation of the real functional relationship is in many cases an impossible task in the light of present knowledge. Lacking the "true form of the law," an empirical relationship is the best approximation. In the discussion which follows, it will be presumed that there is a real functional relationship between the two

variables being considered. It will make no difference whether the relationship can be demonstrated from theoretical grounds or whether the variables have been empirically correlated.¹

The problem of fitting a curve to a set of discordant data resolves itself into two distinct operations. The first is, of necessity, the operation of determining the mathematical form of an equation which can represent the data satisfactorily. The second operation consists in evaluating the constants of the equation as precisely as necessary. If there are theoretical grounds from which to deduce the form of the functional relationship, the problem reduces itself to the second operation only. The choice of an empirical equation to represent a set of data, with no theoretical guidance, is a highly subjective process. Not only is circumspect judgment necessary but also, in many cases, success depends upon the exercise of a great deal of ingenuity. The operation of choosing the form of the curve cannot be reduced to any simple set of rules or steps. The second operation, namely, the evaluation of the constants, is much more straightforward and less subject to the judgment of the computer.

CHOICE OF EQUATION

In this section we shall discuss some of the more frequently used procedures for choosing a functional form to represent a set of data involving two variables x and y . It is almost invariably the custom to consider the independent variable to have been correctly fixed at certain successive values and to regard the dependent variable as carrying the burden of discordancy.

Finite Differences. If an equation of the polynomial type will serve to represent a set of data, the degree of the polynomial may be determined by the method of taking differences. It will be recalled from Chapter V that the column of differences which is constant gives the degree of the polynomial which can represent the function U . This is equally true for differences with either a constant or a variable increment in the independent variable. The method of determining the degree of the polynomial by taking differences will succeed only if the data are fairly concordant.

¹ Students of science who desire an insight into the idea of functional relationship will derive great benefit from a brief study of correlation. One of the books on statistics listed in the Appendix may be consulted.

If differences of very discordant data are taken, the irregularities due to large positive and negative discrepancies may be of the same order of magnitude as the differences themselves and no order or regularity can be discerned.

Occasionally the failure to find a column of constant differences will give a clue as to the form of the functional relationship. Consider the following pairs of values and their differences

x	U_x	Δ	Δ^2	Δ^3	Δ^4
1	8				
		6			
2	14		6		
		12		6	
3	26		12		6
		24		12	
4	50		24		12
		48		24	
5	98		48		
		96			
6	194				

There appears to be no tendency for a column of differences to approach constancy. Another kind of order, however, is apparent. Successive differences of U_1 are 6; likewise successive differences of U_2 are 12; and so on. Furthermore, the values in each of the difference columns are in geometric progression with a constant ratio of 2. This indicates that the function contains 2^x as one term since there are integral values of the independent variable. We should postulate as a general function for the preceding table the formula

$$U_x = b + ma^x \quad (1)$$

Actually the table was made from

$$U_x = 2 + (3)2^x$$

In practice it is extremely unlikely that a set of data would ever show such a simple relationship as that portrayed in this abstract example. Even in this simple case, the recognition of 2^x as a term of the formula was due more to intuition than to mathematical reasoning. In Chapter V we considered only the more elementary aspects of the calculus of finite differences. By a more advanced study of this calculus we can gain a sufficiently powerful method of analysis to deal with more complicated cases. For instance, the constants of equation (1) might have values other than simple integral ones, or the functional relation might have had more terms, as in:

$$U_x = b + cx + dx^2 + ma^{nx} \quad (2)$$

The processes of differencing and summation are capable of yielding generalizations by means of which these complicated functions can be treated analytically. This example has been introduced solely to indicate what must be the line of attack when the functional relationship is too complex for recognition by one of the commonly used methods available in the usual textbooks.

Graphical Method. The given set of data may be plotted on rectangular coordinate paper following the customary procedure of laying off the values of the dependent variable as ordinates and the corresponding values of the independent variable as abscissas. If the relationship is one of the general linear form

$$y = b + mx \quad (3)$$

it will be possible to draw a straight line through the points on the graph. It will be safer to say a straight line can be drawn among the points, for it is only with abstract or idealized data that there will be no deviations. Properly speaking, a point cannot represent a measured value. The only correct representation of a given datum is an ellipse, one axis of which gives the precision of the dependent variable and the other the precision of the independent variable. In accordance with our previous assumption that the discordancy is entirely attributed to the dependent variable, the ellipse will reduce to a vertical line. The length of the line will be such that its extremes give the locations of $a - r_a$ and $a + r_a$.

The accepted conventions for the preparation of graphs of chemical data dictate that the mean value of the dependent variable be denoted by a point and surrounded by a circle the radius of which is the probable error of that value. The circles of precision are, of course, constructed exactly to the scale of the graph.

When it is readily apparent that a smooth curved line is required to pass through the points, it is quite possible that a suitable substitution may reveal that the proper relationship is of a type which can be reduced to a linear equation in the substituted variable. The more frequently used substitutions are the following:

(1) *Plot log x against y.* The logs of the values of the independent variable may be plotted as abscissas against the values of

y as ordinates. If a linear graph results, the data may be represented by an equation of the type

$$y = b + m \log x \quad (4)$$

Equation (4) is seen to be linear if $\log x$ is treated as the independent variable.

(2) *Plot x against $\log y$.* If a plot of the values of x against the corresponding values of $\log y$ results in a straight line, the data may be represented by an equation of the general form

$$y = ka^{nx} \quad (5)$$

Taking logs of both sides,

$$\log y = \log k + (n \log a) x \quad (6)$$

Equation (6) is a linear equation in x and $\log y$.

(3) *Plot $\log x$ against $\log y$.* If a linear relationship is revealed by plotting $\log x$ against $\log y$, the type of formula indicated is

$$y = kx^n \quad (7)$$

Taking logs of both sides of equation (7), we get

$$\log y = \log k + n \log x \quad (8)$$

Equation (8) is linear in the variables $\log x$ and $\log y$.

(4) *Plot x^n against y .* Quite often the data from chemical measurements can be represented by a formula in which the independent variable is raised to a power higher than first. The second power is by far the most common, but occasionally third- and fourth-power functions are required. If a plot of x^n ($n = 2, 3, 4, \dots$) against y results in a straight-line graph, the data can be represented by

$$y = b + mx^n \quad (9)$$

Equation (9) is a linear function, if x^n is treated as the independent variable.

(5) *Plot $\frac{1}{x^n}$ against y .* Plotting $\frac{1}{x^n}$ against y is merely a variant of the method immediately preceding, since $\frac{1}{x^n}$ is the same as x^{-n} . Here, the values of n that may be tried are 1, 2, 3, $\dots$. The frequencies of occurrence are roughly in the order given, and they

should be tried in that order. If a plot of $\frac{1}{x^n}$ against y gives a linear graph, the data can be represented by

$$y = b + \frac{m}{x^n} \quad (10)$$

Equation (10) is a linear equation in $\frac{1}{x^n}$.

(6) *Plot $\sqrt[n]{x}$ against y .* About the only roots of the values of the independent variable that need be tried are the square root and the cube root. If a plot of $\sqrt[n]{x}$ against y yields a straight line, the equation which can represent the data is of the form

$$y = b + mx^{\frac{1}{n}} \quad (11)$$

which is linear in $x^{\frac{1}{n}}$.

(7) *Other possibilities.* The six methods just given will suffice for most of the cases encountered in chemical measurements. It is only remotely possible that such procedures as plotting trigonometric functions of one or the other variable might reveal a linear relationship. There are other possible variants of the methods given, such as

$$x \text{ against } \frac{1}{y},$$

$$\frac{1}{x} \text{ against } \frac{1}{y},$$

$$\sqrt[n]{x} \text{ against } \sqrt[k]{y}.$$

In addition to variants, there are possible combinations of the methods. For example, a plot of $\log x$ against $\log \log y$ giving a straight line would indicate a function like

$$y = a^{mx^n} \quad (12)$$

Smoothing of Data. If a set of data has failed to show any regularity in the process of taking differences, and if the graphical methods have failed to indicate a suitable functional form, the procedure of smoothing the data may be applied. Analytical methods² of smoothing, or graduation, as it is often called, have

² See Whittaker and Robinson: *Calculus of Observations*, Chapter XI, Blackie and Son, London, 1924.

been developed. Some of these methods have a probability basis, and some have no basis. The only method of smoothing we shall consider is the graphical one. After the data have been plotted, a spline of hard rubber or other plastic material can be bent to serve as a guide in drawing a curve. The smooth continuous curve which appears to represent the trend of the data to the highest possible degree is then drawn. From the smooth curve, values of y are taken, preferably with a constant increment in the values of x . The values of y are tabulated, and differences then taken. There should be sufficient concordance now to indicate by a column of constant differences the degree of a polynomial to represent the data. To be on the safe side and make allowance for the low precision of the graphical method (about 0.25 per cent), a polynomial of one more term than the differences indicate may be taken.

Smoothing of data cannot be recommended as a general procedure. It should be pointed out that, in the process of smoothing, the correct functional relationship may be obscured by the arbitrary choice of a smooth curve. Furthermore, slight but significant discontinuities or "breaks" may be overlooked in smoothed data. In so far as these criticisms apply to the promiscuous practice of smoothing as a first resort, they are entirely justified. If, however, smoothing is used as a last resort after all other possibilities have been exhausted, there can be no cause for complaint. If a set of data is so discordant as to require smoothing, the experimenter who made the measurements does not deserve the good fortune of discovering the correct functional relationship. Nor should we place much confidence in the location of discontinuities inferred from such data. For highly discordant data a polynomial of one degree higher than that indicated by the smoothed differences is a sufficiently good representation.

Multiplicity of Equations. When measurements have been performed with very great precision over a wide range of values of the independent variable, it often happens that a single empirical equation is not capable of representing the data with sufficient exactness over the entire range. Sometimes a graph of the data may reveal a portion of the curve over which the curvature is much greater than over other portions so that the portion of greater curvature may require a polynomial of higher degree than that required by the other portions. We then resort to the use of

more than one equation, each of which is applicable to a definite, specified part of the range covered by the variation of the independent variable. This procedure amounts to dividing the data into several independent sets. It is, however, almost invariably the custom to allow the ranges to overlap so that independent checks can be made on the values near the arbitrary borders. The *International Critical Tables* abound with examples of multiplicity of empirical equations for one set of data.

Substituted Variables. We have seen in a preceding section how it is possible to reduce several types of equations to a linear form through substitution of variables. For convenience in subsequent considerations we shall recapitulate the results of the six types of substitution discussed and indicate a convenient method of algebraic substitution.

$$(1) \quad y = b + m \log x.$$

$$\text{Let: } \log x = X$$

Then

$$y = b + mX \tag{13}$$

$$(2) \quad y = ka^{nx}.$$

Taking logs

$$\log y = \log k + (n \log a)x$$

$$\text{Let: } \log y = Y$$

$$\log k = B$$

$$n \log a = M$$

Then

$$Y = B + Mx \tag{14}$$

When M has been evaluated, n can be determined by choosing an arbitrary value of a (usually 10, e , 2, etc.). The choice is immaterial since any exponential system with a given base is consistent within itself.

$$(3) \quad y = kx^n.$$

Taking logs

$$\log y = \log k + n \log x$$

$$\text{Let: } \log y = Y$$

$$\log k = B$$

$$n = M$$

$$\log x = X$$

Then

$$Y = B + MX \quad (15)$$

$$(4) \ y = b + mx^n.$$

Let:

$$x^n = X$$

Then

$$y = b + mX \quad (16)$$

$$(5) \ y = b + \frac{m}{x^n}.$$

Let:

$$\frac{1}{x^n} = X$$

Then

$$y = b + mX \quad (17)$$

$$(6) \ y = b + mx^{\frac{1}{n}}.$$

$$\text{Let: } x^{\frac{1}{n}} = X$$

Then

$$y = b + mX \quad (18)$$

These substituted forms can be treated by the algebraic and graphical methods of evaluating the constants of a real linear equation.

EVALUATION OF CONSTANTS

Several methods are in general use for the evaluation of the constants in an equation. We shall discuss the methods that have been found most useful in the measurements of physical chemistry. The order in which they are presented is also the order of increasing difficulty and approximately the order of increasing objectivity. After the equation type has been chosen, the constants of the equation can be evaluated by one of the succeeding methods.

Interpolation Formulas. If the set of data has shown sufficient concordance to allow differencing, a rough evaluation of the constants can be accomplished by setting up the general form of the interpolation formula. If numerical values of the differences are substituted and the equation cleared to give a polynomial in x , the constants so obtained are a rough approximation which will serve to represent the data over the range for which differences are taken.

Example 1. The data given in Example 2 in Chapter V (p. 72) will serve to illustrate the method. The general interpolation formula is

$$\mu_t = \mu_{20} + m\Delta\mu_{20} + \frac{m(m-1)}{2!} \Delta^2\mu_{20} + \frac{m(m-1)(m-2)}{3!} \Delta^3\mu_{20}$$

and

$$m = \frac{x-a}{h} = \frac{t-20}{2} = \frac{t}{2} - 10$$

Substituting the numerical values of the differences

$$\begin{aligned} \mu_t = 1.0050 + \left(\frac{t}{2} - 10\right)(-0.0471) + \frac{\left(\frac{t}{2} - 10\right)\left(\frac{t}{2} - 10 - 1\right)}{2!}(0.0034) \\ + \frac{\left(\frac{t}{2} - 10\right)\left(\frac{t}{2} - 10 - 1\right)\left(\frac{t}{2} - 10 - 2\right)}{3!}(-0.0002) \end{aligned}$$

This equation can be cleared to yield

$$\mu_t = 1.707 - 0.0485t + 0.00070t^2 - 0.0000042t^3$$

This method of clearing an interpolation formula rarely yields values for the constants which are worth the tedium of the process. If the reader will solve this example in detail, the fault will be seen to lie in the small number of significant figures in the differences.

Graphical Method. The constants of a linear equation, or of an equation that is reducible to linear, can be determined within the precision of the graphical method.³ From the graph, the value of the intercept on the y axis can be read to give the value of b in the equation

$$y = b + mx$$

³ Explicit and detailed directions for graphical procedures may be found in Part II of Goodwin: *Precision of Measurements and Graphical Methods*, McGraw-Hill, New York, 1920.

The slope of the line (m) can be evaluated from the lengths of the legs of a right triangle of which the straight line or a part of it is the hypotenuse. The ratio $\frac{\Delta y}{\Delta x}$ gives the value of m provided that the lengths are measured according to Cartesian convention. Fig. 14 illustrates four cases.

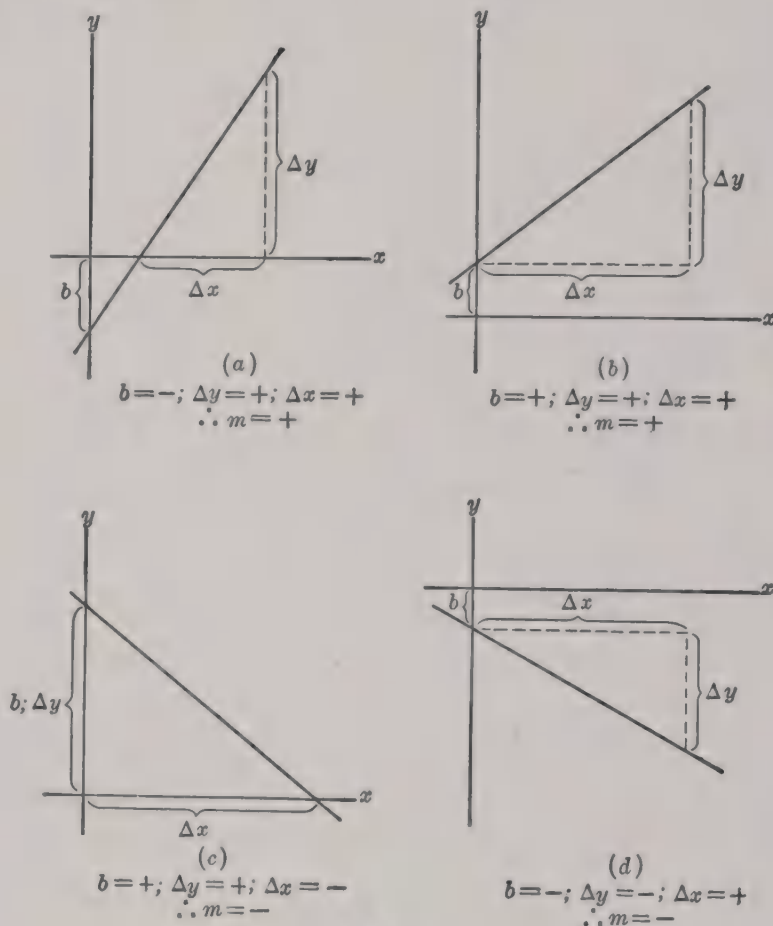


FIG. 14:

The precision of graphical methods has been variously estimated at 0.2 to 0.5 per cent. This means that, if no more than three significant figures in the values of the constants are sought, the graphical evaluation will suffice. In addition to the uncertainty of plotting points, there is the uncertainty involved in the location of the "best" line; hence the overall precision of the

graphical evaluation of the constants will probably be nearer the higher value, 0.5 per cent or 1 part in 200.

Example 2. The following data

x	y	x	y
1	3.0	13	8.0
3	4.0	15	9.0
8	6.0	17	10.0
10	7.0	20	11.0

can be represented by a straight line. By means of a graph, determine the constants of the linear equation

$$y = b + mx$$

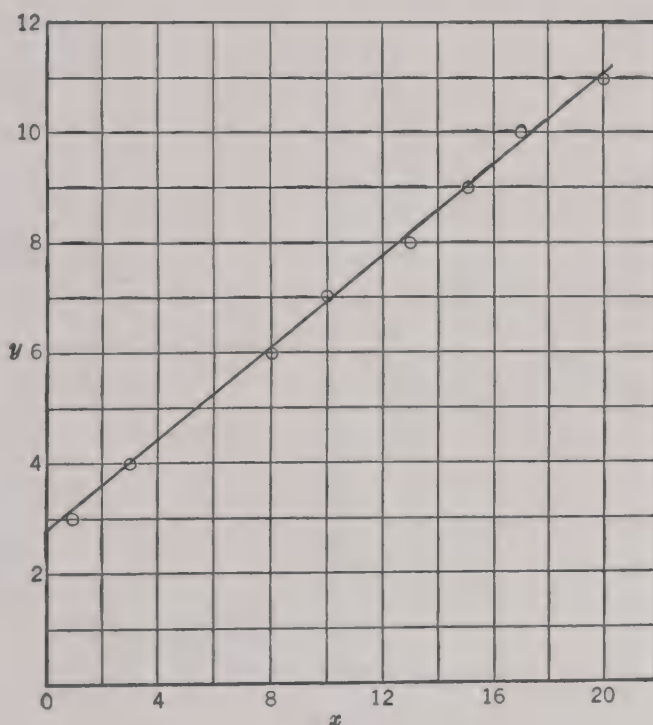


FIG. 15.

Fig. 15 shows a plot of the data with a straight line drawn among the points. The line passes through two points and gives a nearly equal distribution of the remaining points above and below the line. The intercept on the y axis is 2.73; hence

$$b = 2.73$$

To evaluate m , let us read off values of y corresponding to $x = 0$ and $x = 20$. Then

$$m = \frac{\Delta y}{\Delta x} = \frac{11.09 - 2.73}{20.0 - 0} = 0.418$$

Hence the equation for the data is

$$y = 2.73 + 0.418x$$

Simultaneous Equations. If there are k constants in the equation chosen to represent a set of data, the constants can be evaluated if k simultaneous equations are set up. These k equations are formed by choosing, at regularly spaced intervals over the range of the data, k different pairs of corresponding values of x and y and substituting the pairs of numerical values successively in the general equation. Thus, if a linear equation is indicated

$$\left. \begin{aligned} b + m x_p &= y_p \\ b + m x_r &= y_r \end{aligned} \right\} \quad (19)$$

Such a pair of equations when solved simultaneously (determinants may be used) will yield values for b and m .

For highly concordant data, this method will yield numerical values of the constants which may be about as precise as the values obtained by the graphical method. The arbitrary choice of the pairs of values will greatly affect the values of the constants. Any one set of k pairs will give different results from those of any other set of k pairs. It will be recalled from Chapter III that the answers obtained from the solution of simultaneous equations often have fewer significant figures than the data entering the equation. In general, this method of evaluating constants can be classed with the graphical method as to the "goodness" of the results.

Example 3. Let us evaluate the constants of the linear equation for the data of the preceding example by means of simultaneous equations.

x	y	x	y
1	3.0	13	8.0
3	4.0	15	9.0
8	6.0	17	10.0
10	7.0	20	11.0

Suppose the second and seventh pairs of values to be chosen. Then

$$b + 3m = 4.0$$

$$b + 17m = 10.0$$

The solution of these simultaneous equations yields the answers

$$b = 2.72; m = 0.428$$

The equation to represent the data according to this method of evaluating constants is

$$y = 2.72 + 0.428 x$$

It should be apparent that the choice of any other pair of values would give still different answers. Combinatory analysis tells us that, with eight pairs of values, it is possible to choose two at a time in twenty-eight different ways.

$${}_8C_2 = 28$$

We can only conclude that the evaluation of k constants by means of k simultaneous equations is a highly arbitrary process.

Method of Averages. The method of averages is an extension of the immediately preceding method. Instead of forming k equations to evaluate k constants, a series of equations is set up making use of all the n pairs of values in the given data. The series of n equations is then grouped arbitrarily into k sets of equations, each set having approximately the same number of equations. Next, the equations within each set are added together, and the result is k equations which can be solved simultaneously to give numerical values for the k constants. If there were six pairs of values to be fitted with a second-degree equation

$$y = a + bx + cx^2 \quad (20)$$

six equations could be formed

$$\left. \begin{aligned} a + bx_1 + cx_1^2 &= y_1 \\ a + bx_2 + cx_2^2 &= y_2 \\ a + bx_3 + cx_3^2 &= y_3 \\ a + bx_4 + cx_4^2 &= y_4 \\ a + bx_5 + cx_5^2 &= y_5 \\ a + bx_6 + cx_6^2 &= y_6 \end{aligned} \right\} \quad (21)$$

There are fifteen possible ways in which the six equations can be grouped into three sets having two equations each. Therefore, fifteen different solutions are possible for the values of the con-

stants. If we choose to group together the first and fourth, the second and fifth, and the third and sixth, the result is

$$2a + (x_1 + x_4)b + (x_1^2 + x_4^2)c = y_1 + y_4$$

$$2a + (x_2 + x_5)b + (x_2^2 + x_5^2)c = y_2 + y_5$$

$$2a + (x_3 + x_6)b + (x_3^2 + x_6^2)c = y_3 + y_6$$

When the three equations are solved simultaneously, a , b , and c are evaluated. By this method, the constants are functions of the whole set of measurements. Because there are so many possible solutions and because there is no *a priori* basis for choosing which is the best, some degree of uncertainty must remain. Experience has revealed that the best results are obtained when the equations are divided into equal, successive groups. That is, if there are six equations and three constants, the grouping should be

$$\left. \begin{aligned} 2a + (x_1 + x_2)b + (x_1^2 + x_2^2)c &= y_1 + y_2 \\ 2a + (x_3 + x_4)b + (x_3^2 + x_4^2)c &= y_3 + y_4 \\ 2a + (x_5 + x_6)b + (x_5^2 + x_6^2)c &= y_5 + y_6 \end{aligned} \right\} \quad (22)$$

Of the methods considered thus far, this one is the best, and if the data are highly concordant, the constants so determined will give an equation which represents the data very faithfully.

Example 4. To continue the comparison we shall evaluate the constants of the linear equation for the data given in the two preceding examples.

x	y	x	y
1	3.0	13	8.0
3	4.0	15	9.0
8	6.0	17	10.0
10	7.0	20	11.0

There will be eight equations

$$(1) \quad b + m = 3.0$$

$$(5) \quad b + 13m = 8.0$$

$$(2) \quad b + 3m = 4.0$$

$$(6) \quad b + 15m = 9.0$$

$$(3) \quad b + 8m = 6.0$$

$$(7) \quad b + 17m = 10.0$$

$$(4) \quad b + 10m = 7.0$$

$$(8) \quad b + 20m = 11.0$$

Following the rule of dividing the equations into equal successive groups,

the first four equations are placed in one group and the second four in another group. The result is

$b + m = 3.0$	$b + 13m = 8.0$
$b + 3m = 4.0$	$b + 15m = 9.0$
$b + 8m = 6.0$	$b + 17m = 10.0$
$b + 10m = 7.0$	$b + 20m = 11.0$
$4b + 22m = 20.0$	$4b + 65m = 38.0$

The resulting simultaneous equations

$$\begin{aligned} 4b + 22m &= 20.0 \\ 4b + 65m &= 38.0 \end{aligned}$$

yield the values

$$b = 2.70; m = 0.420$$

The linear equation is

$$y = 2.70 + 0.420x$$

Method of Least Squares. It has often been stated that the "method" of least squares for evaluating the constants of an equation is based on the theorems of probability. This is not entirely true. The "principle" of least squares which was enunciated by Legendre is implicitly stated in the Gaussian theory of errors. (See again pages 166-7.) The principle of least squares, therefore, has in common with the Gaussian theory a basis in the theorems of probability. The principle applies to repeated measurements of a physical quantity when all the ponderable variables are fixed. In the evaluation of constants we are dealing with a series of measurements of a quantity corresponding to various fixed values of an independent variable. The two kinds of data are not the same.

The method of least squares in purporting to yield "highly probable" values for the constants of an equation requires two assumptions:

- (1) The fixed values of the independent variable are correct, and hence only the dependent variable is subject to errors of measurement.
- (2) The curve of best fit is the one which makes the sum of the squares of the deviations from the curve a minimum.

Fig. 16 depicts the simplest case, i.e., a linear function. Here, the deviations are magnified out of natural proportion merely for illustrative purposes. Note that as a consequence of assumption (1) the deviations are vertical distances and not perpendicular distances from the curve.

The method which makes the sum of the squares of the deviations a minimum is called the "method" of least squares. It is in reality an intuitive extension of the "principle" of least squares. We might say the method of least squares is a stepchild of the cal-

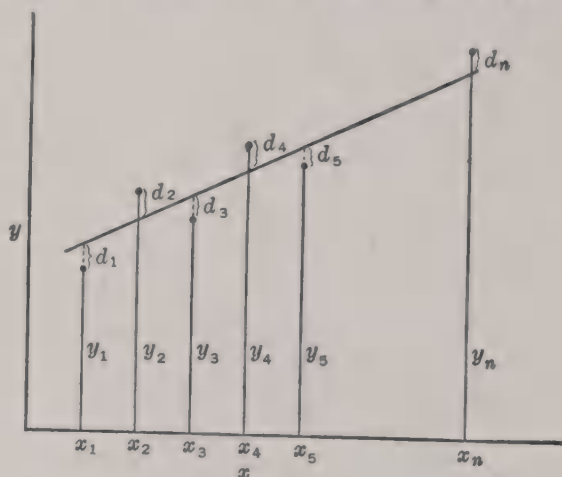


FIG. 16.

culus of probabilities. Having examined the foundations of the method, we shall consider the mathematics involved.

Let us assume that n pairs of values of x and y are to be fitted by the linear equation

$$y = b + mx \quad (23)$$

and we shall denote by y' the values of the dependent variable which could be calculated from values of x when the constants b and m have been evaluated. Then

$$y'_1 = b + mx_1$$

The deviation of the measured value from the curve would be given by

$$d_1 = y_1 - y'_1 = y_1 - (b + mx_1)$$

or

$$d_1 = y_1 - b - mx_1$$

If we let

$$\Sigma d_i^2 = Q$$

then

$$Q = (y_1 - b - mx_1)^2 + (y_2 - b - mx_2)^2 + \dots + (y_n - b - mx_n)^2 \quad (24)$$

In general, if

$$P = f(t, w, z, \dots) \quad (25)$$

the necessary conditions for a minimum are

$$\frac{\partial P}{\partial t} = 0; \quad \frac{\partial P}{\partial w} = 0; \quad \frac{\partial P}{\partial z} = 0; \dots \quad (26)$$

The further conditions need not be applied here.

Since, in equation (24), the values of y_i and x_i have been fixed by measurement, the variables of the equation are b and m . Hence

$$\frac{\partial Q}{\partial b} = 0; \quad \frac{\partial Q}{\partial m} = 0 \quad (27)$$

fix the critical values of b and m for Q to be minimum.

Now

$$\frac{\partial Q}{\partial b} = -2(y_1 - b - mx_1) - 2(y_2 - b - mx_2) - \dots - 2(y_n - b - mx_n)$$

or

$$(y_1 - b - mx_1) + (y_2 - b - mx_2) + \dots + (y_n - b - mx_n) = 0$$

Taking the sum of the terms

$$\Sigma y_i - nb - m \Sigma x_i = 0 \quad (28)$$

Also

$$\begin{aligned} \frac{\partial Q}{\partial m} = & -2x_1(y_1 - b - mx_1) - 2x_2(y_2 - b - mx_2) - \dots \\ & - 2x_n(y_n - b - mx_n) \end{aligned}$$

or

$$\begin{aligned} (x_1 y_1 - b x_1 - m x_1^2) + (x_2 y_2 - b x_2 - m x_2^2) + \dots \\ + (x_n y_n - b x_n - m x_n^2) = 0 \end{aligned}$$

Taking the sum

$$\Sigma x_i y_i - b \Sigma x_i - m \Sigma x_i^2 = 0 \quad (29)$$

The two equations (28) and (29) constitute the *normal equations* for the evaluation of the constants of a linear equation with the two constants b and m .

Without showing the details, we shall state that for a second-degree equation with three constants such as

$$y = a + bx + cx^2 \quad (30)$$

the normal equations are

$$\left. \begin{aligned} a\Sigma x_i^2 + b\Sigma x_i^3 + c\Sigma x_i^4 &= \Sigma x_i^2 y_i \\ a\Sigma x_i + b\Sigma x_i^2 + c\Sigma x_i^3 &= \Sigma x_i y_i \\ an + b\Sigma x_i + c\Sigma x_i^2 &= \Sigma y_i \end{aligned} \right\} \quad (31)$$

Returning to the linear equation, the normal equations can be solved simultaneously to evaluate b and m . Rearranging equations (28) and (29)

$$b\Sigma x_i + m\Sigma x_i^2 = \Sigma x_i y_i$$

$$bn + m\Sigma x_i = \Sigma y_i$$

By the method of determinants

$$b = \frac{\begin{vmatrix} \Sigma x_i y_i & \Sigma x_i^2 \\ \Sigma y_i & \Sigma x_i \end{vmatrix}}{\begin{vmatrix} \Sigma x_i & \Sigma x_i^2 \\ n & \Sigma x_i \end{vmatrix}} = \frac{\Sigma x_i y_i \Sigma x_i - \Sigma y_i \Sigma x_i^2}{(\Sigma x_i)^2 - n \Sigma x_i^2} \quad (32)$$

and

$$m = \frac{\begin{vmatrix} \Sigma x_i & \Sigma x_i y_i \\ n & \Sigma y_i \end{vmatrix}}{\begin{vmatrix} \Sigma x_i & \Sigma x_i^2 \\ n & \Sigma x_i \end{vmatrix}} = \frac{\Sigma x_i \Sigma y_i - n \Sigma x_i y_i}{(\Sigma x_i)^2 - n \Sigma x_i^2} \quad (33)$$

We cannot go so far as to state that the method of least squares will give either the "most probable" values or the "best possible" values of the constants of an equation. But we shall state unreservedly that this method does give the "best fit" of all the methods we have considered. From the standpoint of time required for computation, the method of least squares is undoubtedly the most difficult of the commonly applied procedures. If the experimenter has taken great pains to refine the processes of measurement, a highly objective treatment of the results is not only indicated but completely justified, though the method be difficult.

Before passing on we should note that when a function has been reduced to linear form through substitution the fitting of an

equation by least squares results in a best fit, not of the two variables themselves, but of the substituted variables. Thus if by plotting $\log x$ against $\log y$ a straight line is obtained, indicating the function

$$y = kx^n$$

the reduced linear form is

$$Y = B + MX$$

The curve resulting from the least squares evaluation gives the best fit not in terms of x and y but in terms of $\log x$ and $\log y$. Although the best fit of the substituted variables may not be best for the actual variables, the fit is as good as is readily obtainable by a practicable method.

Example 5. Continuing the comparison we shall use the data of the three previous examples to obtain the constants of the linear equation by the method of least squares. The normal equations are

$$b\Sigma x + m\Sigma x^2 = \Sigma xy$$

$$bn + m\Sigma x = \Sigma y$$

The computation of the required terms can be arranged in tabular form

	x	y	x^2	xy
	1	3.0	1	3.0
	3	4.0	9	12.0
	8	6.0	64	48.0
	10	7.0	100	70.0
	13	8.0	169	104.0
	15	9.0	225	135.0
	17	10.0	289	170.0
	20	11.0	400	220.0
	—	—	—	—
Sum	87	58.0	1257	762.0

The normal equations can be solved when these values are substituted:

$$\Sigma x = 87; \Sigma y = 58.0; \Sigma x^2 = 1,257; \Sigma xy = 762.0; n = 8$$

Expansion of the determinants from the two normal equations gives

$$b = \frac{\Sigma xy \Sigma x - \Sigma y \Sigma x^2}{(\Sigma x)^2 - n \Sigma x^2}$$

and

$$m = \frac{\Sigma x \Sigma y - n \Sigma xy}{(\Sigma x)^2 - n \Sigma x^2}$$

Making the indicated substitutions

$$b = \frac{762 \cdot 87 - 58 \cdot 0 \cdot 1,257}{(87)^2 - 8 \cdot 1257}$$

$$m = \frac{87 \cdot 58 \cdot 0 - 8 \cdot 762}{(87)^2 - 8 \cdot 1257}$$

from which

$$b = 2.66; m = 0.422$$

The least squares equation for the data is

$$y = 2.66 + 0.422 x$$

“Goodness” of Fit. We have spoken of rough, fair, good, and best fits of a curve, but we should like to know “how good.” Reasoning by the same analogies from which the method of least squares was evoked from the principle, we shall make use of the probable error as a measure of goodness of fit. With the constants evaluated it will be possible to compute the values of y'_i from corresponding values of x_i , and hence the values of the deviations may be obtained by

$$d_i = y_i - y'_i$$

The probable error of a single value of the dependent variable is given by Bessel’s formula

$$r = \pm 0.6745 \sqrt{\frac{\sum d_i^2}{n - k}} \quad (34)$$

in which n is the number of pairs of values and k is the number of constants in the fitted equation. The argument for subtracting the number of constants from the number of points is as follows: The evaluation of k constants requires the simultaneous solution of k equations. Now, k pairs of corresponding values substituted in the k equations would give values for the constants such that the resulting curve would pass through those k points and only $n - k$ points would deviate from the curve. Hence, the sum of the squares of the deviations should be divided by the number of deviations. Generalized proofs⁴ show that, for the methods of averages and least squares, the same conclusion is reached. In graphical evaluation, where $k = 2$, the experimenter is likely to be prejudiced in favor of that straight line which passes through at least two points. Lacking more nearly complete information we assume, therefore, that Bessel’s formula (34) gives the probable error of a single value of y by whatever method the constants have

⁴ Whittaker and Robinson, *op. cit.*, pp. 243–5.

been evaluated.⁵ We assume further that the probable error of a single value of y , when consistently employed, serves adequately as a measure of the goodness of fit. Inasmuch as the equation chosen is often only an empirical approximation, we would expect to find that a given set of data can be represented by more than one type of equation. The proper choice when there are two possible equations will be determined by the probable errors. The curve which gives the smaller probable error will be the one chosen.

Example 6. A summary of the results of determining the constants of the linear equation to represent a given set of data is as follows:

$$y = 2.73 + 0.418 x \quad \text{Graphical}$$

$$y = 2.72 + 0.428 x \quad \text{Simultaneous equations}$$

$$y = 2.70 + 0.420 x \quad \text{Averages}$$

$$y = 2.66 + 0.422 x \quad \text{Least squares}$$

Each of the equations has different constants from the others; hence the deviations from the computed values will be different for each equation. To complete our comparison of the methods it will be necessary to compute y'_i for each value of x , by means of each equation. It will then be possible to obtain the deviations and the squares of the deviations for each method. These computations are summarized in the table following:

x	y	Graphical ⁶			Simul. Equa.			Averages			Least Squares		
		y'	$\pm d$	d^2	y'	$\pm d$	d^2	y'	$\pm d$	d^2	y'	$\pm d$	d^2
1	3.0	3.15	0.15	0.0225	3.15	0.15	0.0225	3.12	0.12	0.0144	3.08	0.08	0.0064
3	4.0	4.00	0.00	.0000	4.00	0.00	.0000	3.96	0.04	.0016	3.93	0.07	.0049
8	6.0	6.10	0.10	.0100	6.14	0.14	.0196	6.06	0.06	.0036	6.04	0.04	.0016
10	7.0	6.93	0.07	.0049	7.00	0.00	.0000	6.90	0.10	.0100	6.88	0.12	.0144
13	8.0	8.20	0.20	.0400	8.28	0.28	.0784	8.16	0.16	.0256	8.15	0.15	.0225
15	9.0	9.00	0.00	.0000	9.14	0.14	.0196	9.00	0.00	.0000	8.99	0.01	.0001
17	10.0	9.85	0.15	.0225	10.00	0.00	.0000	9.84	0.16	.0256	9.83	0.17	.0289
20	11.0	11.09	0.09	.0081	11.28	0.28	.0784	11.10	0.10	.0100	11.10	0.10	.0100
Sum		0.1080			0.2185			0.0908			0.0888		

⁵ For methods of evaluating the probable errors of the constants themselves, consult Bond: *Probability and Random Errors*, Chapter VII, Arnold, London, 1935.

⁶ For consistency, the values of y' in the graphical solution were read from the graph, rather than calculated.

An examination of the table reveals that the sum of the squares of the deviations is the least for the method of least squares.

The sums of the squares of the deviations having already been obtained, it will be a relatively simple matter to calculate the probable errors of the four methods of evaluation. Upon making the proper substitutions in the formula

$$r = 0.675 \sqrt{\frac{\Sigma d^2}{n - k}}$$

there is obtained

$$r = 0.675 \sqrt{\frac{0.1080}{8 - 2}} = \pm 0.091 \quad \text{Graphical}$$

$$r = 0.675 \sqrt{\frac{0.2185}{8 - 2}} = \pm 0.13 \quad \text{Simultaneous equations}$$

$$r = 0.675 \sqrt{\frac{0.0908}{8 - 2}} = \pm 0.083 \quad \text{Averages}$$

$$r = 0.675 \sqrt{\frac{0.0888}{8 - 2}} = \pm 0.077 \quad \text{Least squares}$$

In the light of our definition of the probable error as rigorously applied (see p. 169), we should anticipate that as here applied it would give the limit within which 50 per cent of the deviations would fall. By reference to the table of deviations we can test this postulate.

Graphical	50% fall within ± 0.09
Simultaneous equations	37% " " ± 0.13
Averages	37% " " ± 0.08
Least squares	50% " " ± 0.08

Considering the small number of points involved, the agreement is very good.

PROBLEM SET XV

1. The following data give the boiling points of several members of an homologous series of hydrocarbons:

HYDROCARBON	BOILING POINT (°C)
C_4H_{10}	0.6
C_5H_{12}	36.2
C_6H_{14}	69.0
C_7H_{16}	94.8
C_8H_{18}	124.6
C_9H_{20}	150.6
$C_{10}H_{22}$	174.0

(a) Prove that the data can be represented by the equation⁷

$$T = aM^b$$

⁷ Walker: *J. Chem. Soc.*, 65, 193 (1894).

in which T = boiling point on absolute scale.

M = molecular weight.

a and b are empirical constants.

(b) Evaluate the constants a and b graphically.

2. Determine the degree of the polynomial required to represent the relationship between the atomic heat capacity of vanadium and the absolute temperature, as given by Anderson.⁸

$T^\circ \text{K}$	C_p	$T^\circ \text{K}$	C_p
74.3	2.181	120.3	3.851
80.2	2.441	135.7	4.248
91.5	2.871	151.0	4.596
101.6	3.290	185.2	5.143
110.2	3.564	206.7	5.358

3. Determine the functional relationship between x and y in these data.

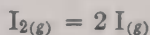
x	y	x	y
0.1	1.0101	0.9	2.2479
0.2	1.0408	1.3	5.4194
0.4	1.1735	1.5	9.4877
0.7	1.6323	1.9	30.6966

4. Moeller and Krauskopf⁹ have determined the densities at 25°C of hydrous lanthanum oxide sols of various concentrations.

La_2O_3 (g per l)	DENSITY (g per cc)
4.04	1.0005
3.64	1.0002
3.23	0.9998
2.83	0.9994
2.42	0.9991
2.02	0.9988
1.62	0.9984
1.21	0.9981
0.81	0.9977
0.40	0.9973
0.00	0.9970

A linear equation can represent density as a function of concentration. Evaluate the constants graphically.

5. The reaction



has been studied by Starck and Bodenstein.¹⁰ The values of the equilibrium constants for this reaction at various temperatures are:

$T^\circ \text{A}$	K	$T^\circ \text{A}$	K
1,073	0.0114	1,373	0.492
1,173	0.0474	1,473	1.23
1,273	0.165		

⁸ *J. Am. Chem. Soc.*, **53**, 565 (1936). ⁹ *J. Phys. Chem.*, **43**, 363 (1939)

¹⁰ *Z. Elektrochem.*, **16**, 961 (1910).

- (a) Determine the form of the equation which can represent the data and
 (b) evaluate the constants by the method of averages.

Hint: Try a plot of T^{-n} against $\log K$.

6. Repeat the solution of the example under "Method of Averages," page 216, changing the grouping of the eight equations as follows: (a) odd-numbered equations in one group and even-numbered in the second group; and (b) (1), (2), (5), (6) in one group and (3), (4), (7), (8) in the second group.

Compare the constants so obtained with the summary on page 223.

7. Measurements of light transmitted by aqueous solutions of $\text{Cu}(\text{NH}_3)_4^{++}$ ions gave the following photocell responses expressed in microamperes. The data are from Yoe and Crumpler.¹¹

C (p.p.m. Cu^{++})	R (μa)	C (p.p.m. Cu^{++})	R (μa)
0	50.0	20	38.0
5	46.8	25	35.3
10	43.7	35	30.8
15	40.9	50	25.1

- (a) Determine the type of equation required to represent these data;
 (b) evaluate the constants by the method of least squares.

8. In the following table are given the wave lengths of the first x-ray spectral line ($K\alpha$) emitted by the indicated elements. The values of N give the order in which the elements are placed in the periodic table of Mendelyev.

ELEMENT	N	λ (cm)
Al	13	8.364×10^{-8}
Si	14	7.142 "
Cl	17	4.750 "
K	19	3.759 "
Ca	20	3.368 "
Ti	22	2.758 "
V	23	2.519 "
Cr	24	2.301 "
Mn	25	2.111 "
Fe	26	1.946 "
Co	27	1.798 "
Ni	28	1.662 "
Cu	29	1.549 "
Zn	30	1.445 "

A plot of N against a substituted variable, $Y = f(\lambda)$, results in a linear graph. Determine the functional relation expressed as an equation. For an interpretation of the result, the reader may consult Moseley, *Phil. Mag.*, **26**, 1024 (1913); *ibid.*, **27**, 703 (1914), from which papers these data were taken.

¹¹ *Ind. Eng. Chem., Anal. Ed.*, **7**, 281 (1935).

APPENDIX

BIBLIOGRAPHY

CLASSIFIED REFERENCES FOR FURTHER STUDY

The scope of this book has not permitted the fullness of treatment of each subject that many readers may desire. To serve as a guide for further study we have prepared the following list of principal topics, arranged approximately in the order of treatment in the book. Under each topic are indicated several numbers which refer to the references listed alphabetically on the succeeding pages. Because some of the textbooks constitute excellent sources of reference on many topics, we have chosen this numerical scheme to eliminate needless duplication.

Historical 2, 8, 9.
Slide Rule 7, 9, 12, 35.
General Algebra 3, 15, 36, 38, 40.
Determinants 15, 29, 36, 38.
Calculus 12, 29, 32, 34, 37, 38, 41, 42.
Interpolation and Extrapolation 1, 5, 29, 30, 34, 37, 39, 47.
Theory of Measurement 10, 19, 26, 29, 33, 43, 47.
Standards of Measurement 16, 20, 31.
Physicochemical Measurements 14, 23, 24, 26, 28, 35, 46.
Statistics 6, 18, 27.
Probability 4, 12, 17, 18, 29, 38, 47, 48.
Theory of Errors 4, 13, 18, 19, 26, 27, 33, 35, 37, 43, 46, 47, 48.
Graphical Methods 12, 19, 37, 43.
Curve Fitting 4, 12, 19, 27, 34, 37, 38, 43, 47.
Tables 11, 16, 21, 22, 25, 44, 45.

1. ARENSON: *Chemical Arithmetic*, Wiley, New York, 1931.
2. BALL: *History of Mathematics*, Second Edition, Macmillan, London, 1893.
3. BOCHER: *Introduction to Advanced Algebra*, Macmillan, New York, 1924.
4. BOND: *Probability and Random Errors*, Arnold, London, 1935.
5. BOOLE: *A Treatise on the Calculus of Finite Differences*, Third Edition, Macmillan, London, 1880.
6. BOWLEY: *Elements of Statistics*, Fifth Edition, King, London, 1926.

7. BRECKENRIDGE: Various self-teaching manuals of the slide rule, published by Keuffel and Esser Co., New York.
8. CAJORI: *History of Mathematics*, Second Edition, Macmillan, New York, 1919.
9. CAJORI: *History of the Slide Rule*, Engineering News Publishing Co., New York, 1909.
10. CAMPBELL: *Measurement and Calculations*, Longmans, Green, London, 1928.
11. CHAPPELL: *Five-Figure Mathematical Tables*, Van Nostrand, London, 1915.
12. DANIELS: *Mathematical Preparation for Physical Chemistry*, McGraw-Hill, New York, 1928.
13. DEMING and BIRGE: "On the Statistical Theory of Errors," *Reviews of Modern Physics*, **6**, 119 (1934).
14. FALES and KENNY: *Inorganic Quantitative Analysis*, new edition. Appleton-Century, New York, 1939.
15. FINE: *College Algebra*, Ginn, Boston, 1905.
16. FOWLE: *Smithsonian Physical Tables*, Eighth Edition, Smithsonian Institution, Washington, 1934.
17. FRY: *Probability and Its Engineering Uses*, Van Nostrand, New York, 1928.
18. GALE and WATKEYS: *Elementary Functions*, Holt, New York, 1920.
19. GOODWIN: *Precision of Measurements and Graphical Methods*, McGraw-Hill, New York, 1920.
20. HALLOCK and WADE: *Evolution of Weights and Measures*, Macmillan, New York, 1906.
21. HODGMAN: *Handbook of Chemistry and Physics*, revised annually, Chemical Rubber Publishing Co., Cleveland.
22. KAYE and LABY: *Physical and Chemical Constants*, Longmans, Green, New York, 1921.
23. KOLTHOFF and SANDELL: *Textbook of Quantitative Inorganic Analysis*, Macmillan, New York, 1936.
24. LACEY: *Instrumental Methods of Chemical Analysis*, Macmillan, New York, 1924.
25. LANGE: *Handbook of Chemistry*, Second Edition, Handbook Publishers, Sandusky, Ohio, 1937.
26. LIVINGSTON: *Physico-Chemical Experiments*, Macmillan, New York, 1939.
27. LOVITT and HOLTZCLAW: *Statistics*, Prentice-Hall, New York, 1929.
28. LUNDELL and HOFFMAN: *Outlines of Chemical Analysis*, Wiley, New York, 1938.
29. MELLOR: *Higher Mathematics for Students of Physics and Chemistry*, Longmans, Green, New York, 1919.

30. MILNE-THOMSON: *The Calculus of Finite Differences*, Macmillan, New York, 1933.
31. National Bureau of Standards *Miscellaneous Publication M 121*, 1936.
32. OSBORNE: *Differential and Integral Calculus*, Heath, Boston, 1908.
33. PALMER: *The Theory of Measurements*, McGraw-Hill, New York, 1912.
34. PARTINGTON: *Higher Mathematics for Chemical Students*, Methuen, London, 1911.
35. REILLY, RAE, and WHEELER: *Physico-Chemical Methods*, Van Nostrand, New York, 1925.
36. SHAW: *Freshman Algebra*, Crowell, New York, 1929.
37. SHERWOOD and REED: *Applied Mathematics in Chemical Engineering*, McGraw-Hill, New York, 1939.
38. SOKOLNIKOFF and SOKOLNIKOFF: *Higher Mathematics for Engineers and Physicists*, McGraw-Hill, New York, 1934.
39. STEFFENSEN: *Interpolation*, Williams and Wilkins, Baltimore, 1927.
40. THOMPSON, J. E.: *Algebra for the Practical Man*, Van Nostrand, New York, 1931.
41. THOMPSON, J. E.: *The Calculus for the Practical Man*. Van Nostrand, New York, 1931.
42. THOMPSON, S. P.: *Calculus Made Easy*, Second Edition, Macmillan, New York, 1919.
43. TUTTLE and SATTERLY: *The Theory of Measurements*, Longmans, Green, New York, 1925.
44. VON VEGA and FISCHER: *Logarithmic Tables* (Seven-place), Weidmanns, Berlin, 1867.
45. WASHBURN (Editor): *International Critical Tables*, eight volumes, McGraw-Hill, New York, 1926-1933.
46. WHITEHEAD: *Design and Use of Instruments*, Macmillan, New York, 1934.
47. WHITTAKER and ROBINSON: *Calculus of Observations*, Blackie and Son, London, 1924.
48. WOODWARD: *Probability and Theory of Errors*, Wiley, New York, 1906.

TABLE I
FOUR-PLACE LOGS

N	0	1	2	3	4	5	6	7	8	9
1.0	.0000	.0043	.0086	.0128	.0170	.0212	.0253	.0294	.0334	.0374
1.1	.0414	.0453	.0492	.0531	.0569	.0607	.0645	.0682	.0719	.0755
1.2	.0792	.0828	.0864	.0899	.0934	.0969	.1004	.1038	.1072	.1106
1.3	.1139	.1173	.1206	.1239	.1271	.1303	.1335	.1367	.1399	.1430
1.4	.1461	.1492	.1523	.1553	.1584	.1614	.1644	.1673	.1703	.1732
1.5	.1761	.1790	.1818	.1847	.1875	.1903	.1931	.1959	.1987	.2014
1.6	.2041	.2068	.2095	.2122	.2148	.2175	.2201	.2227	.2253	.2279
1.7	.2304	.2330	.2355	.2380	.2405	.2430	.2455	.2480	.2504	.2529
1.8	.2553	.2577	.2601	.2625	.2648	.2672	.2695	.2718	.2742	.2765
1.9	.2788	.2810	.2833	.2856	.2878	.2900	.2923	.2945	.2967	.2989
2.0	.3010	.3032	.3054	.3075	.3096	.3118	.3139	.3160	.3181	.3201
2.1	.3222	.3243	.3263	.3284	.3304	.3324	.3345	.3365	.3385	.3404
2.2	.3424	.3444	.3464	.3483	.3502	.3522	.3541	.3560	.3579	.3598
2.3	.3617	.3636	.3655	.3674	.3692	.3711	.3729	.3747	.3766	.3784
2.4	.3802	.3820	.3838	.3856	.3874	.3892	.3909	.3927	.3945	.3962
2.5	.3979	.3997	.4014	.4031	.4048	.4065	.4082	.4099	.4116	.4133
2.6	.4150	.4166	.4183	.4200	.4216	.4232	.4249	.4265	.4281	.4298
2.7	.4314	.4330	.4346	.4362	.4378	.4393	.4409	.4425	.4440	.4456
2.8	.4472	.4487	.4502	.4518	.4533	.4548	.4564	.4579	.4594	.4609
2.9	.4624	.4639	.4654	.4669	.4683	.4698	.4713	.4728	.4742	.4757
3.0	.4771	.4786	.4800	.4814	.4829	.4843	.4857	.4871	.4886	.4900
3.1	.4914	.4928	.4942	.4955	.4969	.4983	.4997	.5011	.5024	.5038
3.2	.5051	.5065	.5079	.5092	.5105	.5119	.5132	.5145	.5159	.5172
3.3	.5185	.5198	.5211	.5224	.5237	.5250	.5263	.5276	.5289	.5302
3.4	.5315	.5328	.5340	.5353	.5366	.5378	.5391	.5403	.5416	.5428
3.5	.5441	.5453	.5465	.5478	.5490	.5502	.5514	.5527	.5539	.5551
3.6	.5563	.5575	.5587	.5599	.5611	.5623	.5635	.5647	.5658	.5670
3.7	.5682	.5694	.5705	.5717	.5729	.5740	.5752	.5763	.5775	.5786
3.8	.5798	.5809	.5821	.5832	.5843	.5855	.5866	.5877	.5888	.5899
3.9	.5911	.5922	.5933	.5944	.5955	.5966	.5977	.5988	.5999	.6010
4.0	.6021	.6031	.6042	.6053	.6064	.6075	.6085	.6096	.6107	.6117
4.1	.6128	.6138	.6149	.6160	.6170	.6180	.6191	.6201	.6212	.6222
4.2	.6232	.6243	.6253	.6263	.6274	.6284	.6294	.6304	.6314	.6325
4.3	.6335	.6345	.6355	.6365	.6375	.6385	.6395	.6405	.6415	.6425
4.4	.6435	.6444	.6454	.6464	.6474	.6484	.6493	.6503	.6513	.6522
4.5	.6532	.6542	.6551	.6561	.6571	.6580	.6590	.6599	.6609	.6618
4.6	.6628	.6637	.6646	.6656	.6665	.6675	.6684	.6693	.6702	.6712
4.7	.6721	.6730	.6739	.6749	.6758	.6767	.6776	.6785	.6794	.6803
4.8	.6812	.6821	.6830	.6839	.6848	.6857	.6866	.6875	.6884	.6893
4.9	.6902	.6911	.6920	.6928	.6937	.6946	.6955	.6964	.6972	.6981
5.0	.6990	.6998	.7007	.7016	.7024	.7033	.7042	.7050	.7059	.7067
5.1	.7076	.7084	.7093	.7101	.7110	.7118	.7126	.7135	.7143	.7152
5.2	.7160	.7168	.7177	.7185	.7193	.7202	.7210	.7218	.7226	.7235
5.3	.7243	.7251	.7259	.7267	.7275	.7284	.7292	.7300	.7308	.7316
5.4	.7324	.7332	.7340	.7348	.7356	.7364	.7372	.7380	.7388	.7396
N	0	1	2	3	4	5	6	7	8	9

TABLE I—(Continued)

FOUR-PLACE LOGS

N	0	1	2	3	4	5	6	7	8	9
5.5	.7404	.7412	.7419	.7427	.7435	.7443	.7451	.7459	.7466	.7474
5.6	.7482	.7490	.7497	.7505	.7513	.7520	.7528	.7536	.7543	.7551
5.7	.7559	.7566	.7574	.7582	.7589	.7597	.7604	.7612	.7619	.7627
5.8	.7634	.7642	.7649	.7657	.7664	.7672	.7679	.7686	.7694	.7701
5.9	.7709	.7716	.7723	.7731	.7738	.7745	.7752	.7760	.7767	.7774
6.0	.7782	.7789	.7796	.7803	.7810	.7818	.7825	.7832	.7839	.7846
6.1	.7853	.7860	.7868	.7875	.7882	.7889	.7896	.7903	.7910	.7917
6.2	.7924	.7931	.7938	.7945	.7952	.7959	.7966	.7973	.7980	.7987
6.3	.7993	.8000	.8007	.8014	.8021	.8028	.8035	.8041	.8048	.8055
6.4	.8062	.8069	.8075	.8082	.8089	.8096	.8102	.8109	.8116	.8122
6.5	.8129	.8136	.8142	.8149	.8156	.8162	.8169	.8176	.8182	.8189
6.6	.8195	.8202	.8209	.8215	.8222	.8228	.8235	.8241	.8248	.8254
6.7	.8261	.8267	.8274	.8280	.8287	.8293	.8299	.8306	.8312	.8319
6.8	.8325	.8331	.8338	.8344	.8351	.8357	.8363	.8370	.8376	.8382
6.9	.8388	.8395	.8401	.8407	.8414	.8420	.8426	.8432	.8439	.8445
7.0	.8451	.8457	.8463	.8470	.8476	.8482	.8488	.8494	.8500	.8506
7.1	.8513	.8519	.8525	.8531	.8537	.8543	.8549	.8555	.8561	.8567
7.2	.8573	.8579	.8585	.8591	.8597	.8603	.8609	.8615	.8621	.8627
7.3	.8633	.8639	.8645	.8651	.8657	.8663	.8669	.8675	.8681	.8686
7.4	.8692	.8698	.8704	.8710	.8716	.8722	.8727	.8733	.8739	.8745
7.5	.8751	.8756	.8762	.8768	.8774	.8779	.8785	.8791	.8797	.8802
7.6	.8808	.8814	.8820	.8825	.8831	.8837	.8842	.8848	.8854	.8859
7.7	.8865	.8871	.8876	.8882	.8887	.8893	.8899	.8904	.8910	.8915
7.8	.8921	.8927	.8932	.8938	.8943	.8949	.8954	.8960	.8965	.8971
7.9	.8976	.8982	.8987	.8993	.8998	.9004	.9009	.9015	.9020	.9025
8.0	.9031	.9036	.9042	.9047	.9053	.9058	.9063	.9069	.9074	.9079
8.1	.9085	.9090	.9096	.9101	.9106	.9112	.9117	.9122	.9128	.9133
8.2	.9138	.9143	.9149	.9154	.9159	.9165	.9170	.9175	.9180	.9186
8.3	.9191	.9196	.9201	.9206	.9212	.9217	.9222	.9227	.9232	.9238
8.4	.9243	.9248	.9253	.9258	.9263	.9269	.9274	.9279	.9284	.9289
8.5	.9294	.9299	.9304	.9309	.9315	.9320	.9325	.9330	.9335	.9340
8.6	.9345	.9350	.9355	.9360	.9365	.9370	.9375	.9380	.9385	.9390
8.7	.9395	.9400	.9405	.9410	.9415	.9420	.9425	.9430	.9435	.9440
8.8	.9445	.9450	.9455	.9460	.9465	.9469	.9474	.9479	.9484	.9489
8.9	.9494	.9499	.9504	.9509	.9513	.9518	.9523	.9528	.9533	.9538
9.0	.9542	.9547	.9552	.9557	.9562	.9566	.9571	.9576	.9581	.9586
9.1	.9590	.9595	.9600	.9605	.9609	.9614	.9619	.9624	.9628	.9633
9.2	.9638	.9643	.9647	.9652	.9657	.9661	.9666	.9671	.9675	.9680
9.3	.9685	.9689	.9694	.9699	.9703	.9708	.9713	.9717	.9722	.9727
9.4	.9731	.9736	.9741	.9745	.9750	.9754	.9759	.9763	.9768	.9773
9.5	.9777	.9782	.9786	.9791	.9795	.9800	.9805	.9809	.9814	.9818
9.6	.9823	.9827	.9832	.9836	.9841	.9845	.9850	.9854	.9859	.9863
9.7	.9868	.9872	.9877	.9881	.9886	.9890	.9894	.9899	.9903	.9908
9.8	.9912	.9917	.9921	.9926	.9930	.9934	.9939	.9943	.9948	.9952
9.9	.9956	.9961	.9965	.9969	.9974	.9978	.9983	.9987	.9991	.9996
N	0	1	2	3	4	5	6	7	8	9

TABLE II
ERROR INTEGRAL *

$$\Phi(t) = \frac{1}{\sqrt{\pi}} \int_{-t}^t e^{-t^2} dt$$

t	0	1	2	3	4	5	6	7	8	9
0.0		.01128	.02256	.03384	.04511	.05637	.06762	.07886	.09008	.10128
1	.11246	.12362	.13476	.14587	.15695	.16800	.17901	.18999	.20094	.21184
2	.22270	.23352	.24430	.25502	.26570	.27633	.28690	.29742	.30788	.31828
3	.32863	.33891	.34913	.35928	.36936	.37938	.38933	.39921	.40901	.41874
4	.42839	.43797	.44747	.45689	.46623	.47548	.48466	.49375	.50275	.51167
5	.52050	.52924	.53790	.54646	.55494	.56332	.57162	.57982	.58792	.59594
6	.60386	.61168	.61941	.62705	.63459	.64203	.64938	.65663	.66378	.67084
7	.67780	.68467	.69143	.69810	.70468	.71116	.71754	.72382	.73001	.73610
8	.74210	.74800	.75381	.75952	.76514	.77067	.77610	.78144	.78669	.79184
9	.79691	.80188	.80677	.81156	.81627	.82089	.82542	.82987	.83423	.83851
1.0	.84270	.84681	.85084	.85478	.85865	.86244	.86614	.86977	.87333	.87680
1	.88021	.88353	.88679	.88997	.89308	.89612	.89910	.90200	.90484	.90761
2	.91031	.91296	.91553	.91805	.92051	.92290	.92524	.92751	.92973	.93190
3	.93401	.93606	.93807	.94002	.94191	.94376	.94556	.94731	.94902	.95067
4	.95229	.95385	.95538	.95686	.95830	.95970	.96105	.96237	.96365	.96490
5	.96611	.96728	.96841	.96952	.97059	.97162	.97263	.97360	.97455	.97546
6	.97635	.97721	.97804	.97884	.97962	.98038	.98110	.98181	.98249	.98315
7	.98379	.98441	.98500	.98558	.98613	.98667	.98719	.98769	.98817	.98864
8	.98909	.98952	.98994	.99035	.99074	.99111	.99147	.99182	.99216	.99248
9	.99279	.99309	.99338	.99366	.99392	.99418	.99443	.99466	.99489	.99511
2.0	.99532	.99552	.99572	.99591	.99609	.99626	.99642	.99658	.99673	.99688
1	.99702	.99715	.99728	.99741	.99753	.99764	.99775	.99785	.99795	.99805
2	.99814	.99822	.99831	.99839	.99846	.99854	.99861	.99867	.99874	.99880
3	.99886	.99891	.99897	.99902	.99906	.99911	.99915	.99920	.99924	.99928
4	.99931	.99935	.99938	.99941	.99944	.99947	.99950	.99952	.99955	.99957
5	.99959	.99961	.99963	.99965	.99967	.99969	.99971	.99972	.99974	.99975
6	.99976	.99978	.99979	.99980	.99981	.99982	.99983	.99984	.99985	.99986
7	.99987	.99987	.99988	.99989	.99989	.99990	.99991	.99991	.99992	.99992
8	.99992	.99993	.99993	.99994	.99994	.99994	.99995	.99995	.99995	.99996
9	.99996	.99996	.99996	.99997	.99997	.99997	.99997	.99997	.99997	.99998
3	.999 98									
4	.999 999 984									
5	.999 999 999 9984									

* Burgess, *Trans. Roy. Soc. (Edinburgh)*, **39**, 257 (1900).

TABLE III
INTERNATIONAL ATOMIC WEIGHTS
1939

PUBLISHED BY THE JOURNAL OF THE AMERICAN CHEMICAL SOCIETY

	Sym- bol	Atomic Number	Atomic Weight		Sym- bol	Atomic Number	Atomic Weight
Aluminum	Al	13	26.97	Molybdenum	Mo	42	95.95
Antimony	Sb	51	121.76	Neodymium	Nd	60	144.27
Argon	A	18	39.944	Neon	Ne	10	20.183
Arsenic	As	33	74.91	Nickel	Ni	28	58.69
Barium	Ba	56	137.36	Nitrogen	N	7	14.008
Beryllium	Be	4	9.02	Osmium	Os	76	190.2
Bismuth	Bi	83	209.00	Oxygen	O	8	16.0000
Boron	B	5	10.82	Palladium	Pd	46	106.7
Bromine	Br	35	79.916	Phosphorus	P	15	30.98
Cadmium	Cd	48	112.41	Platinum	Pt	78	195.23
Calcium	Ca	20	40.08	Potassium	K	19	39.096
Carbon	C	6	12.010	Praseodymium	Pr	59	140.92
Cerium	Ce	58	140.13	Protactinium	Pa	91	231
Cesium	Cs	55	132.91	Radium	Ra	88	226.05
Chlorine	Cl	17	35.457	Radon	Rn	86	222
Chromium	Cr	24	52.01	Rhenium	Re	75	186.31
Cobalt	Co	27	58.94	Rhodium	Rh	45	102.91
Columbium	Cb	41	92.91	Rubidium	Rb	37	85.48
Copper	Cu	29	63.57	Ruthenium	Ru	44	101.7
Dysprosium	Dy	66	162.46	Samarium	Sm	62	150.43
Erbium	Er	68	167.2	Scandium	Sc	21	45.10
Europium	Eu	63	152.0	Selenium	Se	34	78.96
Fluorine	F	9	19.00	Silicon	Si	14	28.06
Gadolinium	Gd	64	156.9	Silver	Ag	47	107.880
Gallium	Ga	31	69.72	Sodium	Na	11	22.997
Germanium	Ge	32	72.60	Strontium	Sr	38	87.63
Gold	Au	79	197.2	Sulfur	S	16	32.06
Hafnium	Hf	72	178.6	Tantalum	Ta	73	180.88
Helium	He	2	4.003	Tellurium	Te	52	127.61
Holmium	Ho	67	163.5	Terbium	Tb	65	159.2
Hydrogen	H	1	1.0081	Thallium	Tl	81	204.39
Indium	In	49	114.76	Thorium	Th	90	232.12
Iodine	I	53	126.92	Thulium	Tm	69	169.4
Iridium	Ir	77	193.1	Tin	Sn	50	118.70
Iron	Fe	26	55.84	Titanium	Ti	22	47.90
Krypton	Kr	36	83.7	Tungsten	W	74	183.92
Lanthanum	La	57	138.92	Uranium	U	92	238.07
Lead	Pb	82	207.21	Vanadium	V	23	50.95
Lithium	Li	3	6.940	Xenon	Xe	54	131.3
Lutecium	Lu	71	175.0	Ytterbium	Yb	70	173.04
Magnesium	Mg	12	24.32	Yttrium	Y	39	88.92
Manganese	Mn	25	54.93	Zinc	Zn	30	65.38
Mercury	Hg	80	200.61	Zirconium	Zr	40	91.22

TABLE IV
ANSWERS TO PROBLEMS

PROBLEM SET I

- | | |
|---|---|
| 1. (a) 6.75×10^4
(b) 3.41×10^2
(c) 9.65217×10^{12}
(d) 7×10^{-10}
(e) 3.026×10^{-3}
(f) 2.01×10^{-2} | 2. (a) 40,000,000,000,000
(b) 0.000 003 67
(c) 0.4028
(d) 140,000
(e) 3.9
(f) 7314.5 |
| 3. (a) 2.30×10^{-5}
(b) 7.652×10^2
(c) 8.93×10^0
(d) 7.4×10^7
(e) 1.001×10^{-4}
(f) 5×10^{-2} | 4. (a) 0.000 128
(b) 0.040
(c) 85,031,000
(d) 0.000 000 000 2
(e) 0.0312
(f) 1.03 |

PROBLEM SET II

- | | |
|--|--|
| 1. (a) 9.99×10^{-7}
(b) 2.05×10^{14}
(c) 8.4×10^{16}
(d) 3.875×10^{-12}
(e) 6.45
(f) 8.06×10^{23} | 2. (a) 3.0×10^{-21}
(b) 6.0×10^3
(c) 1.10×10^{-10}
(d) 9.0×10^4
(e) 3.0×10^{30}
(f) 1.36×10^{-12} |
| 3. (a) 1.62×10^6
(b) 2.162×10^{-5}
(c) 3.113×10^{19}
(d) 4.2×10^{-9}
(e) 1.9×10^3
(f) 2.6×10^{15} | 4. (a) 4.2×10^7
(b) 4.1×10^{-2}
(c) 3.6×10^2 |

PROBLEM SET III

- | | |
|---|---|
| 1. (a) 3.6×10^{-7}
(b) 1.728×10^{21}
(c) 1.6×10^{49}
(d) 2.16×10^{-10} | 2. (a) 8.0×10^{-6}
(b) 4.0×10^3
(c) 9.0×10^6
(d) 6.0×10^{-7} |
| 3. (a) 4.3×10^{-4}
(b) 6.6×10^6
(c) 1.1×10^{-7}
(d) 8.98×10^{-3} | 4. (a) 2.5×10^7
(b) 7.0×10^{-7}
(c) 9.43×10^3 |

PROBLEM SET IV

2. (a) 3.49
(b) 11.10
(c) 5.365
3. (a) 4×10^{-1}
(b) 2.5×10^{-5}
(c) 2×10^{-8}
4. (a) 35.96
(b) 7.17
(c) 9.2
5. (a) 0.790
(b) 3.04×10^{-2}

PROBLEM SET V

1. (a) 209
(b) 93.8
(c) 104
(d) 0.916
(e) 154
2. (a) 0.0845
(b) 1.47
(c) 7.45
(d) 27.6
3. (a) 5.19
(b) 2.09
(c) 2.90
(d) 2.04
4. (a) 11.66
(b) 104.0
(c) 1.19×10^{-5}
(d) 4.07

PROBLEM SET VI

1. (a) 275 kg
(b) 271 mm
(c) 2 cents
(d) 6.4 ml
(e) 17 g
2. (a) 1.08
(b) 0.862
(c) 27.1
(d) 8.97×10^{-16}
(e) 3.14
3. 8.69×10^{-3}
Slide Rule
4. (a) 18.6
(b) 1.70
5. 13.9
6. 0.0756
7. 4.1×10^{-12}
8. $5 \frac{29}{32}$
9. (a) 5.07
(b) $\pm 1 \times 10^{-35}$
10. 8.6

PROBLEM SET VII

1. (a) 6.603
(b) 3.832
2. (a) -0.782
(b) 0.6012
(c) 0.89
(d) 2.356
3. 0.00746 g NH_4Cl
0.05617 g KCl
4. $x = 0.5$
 $y = 3$
5. $x = 2$
 $y = 5$
 $z = 1$
6. 46.1% CaCO_3
53.8% MgCO_3

7. 0.0332 g HNO_3
0.0703 g HNO_2

9. (a) $1\frac{1}{2}$; $-\frac{2}{3}$
(b) 3; -2.5
(c) 8.539×10^{-2} ; 1.32×10^{-3}

10. 2.1

PROBLEM SET VIII

1. 7.0×10^{-8}

2. 0.0022

3. 0.125

4. 473.0

5. 1.40

6. 475.89

PROBLEM SET IX

2. 0.2427

3. 13.5581

0.2482

13.5409

4. 64.7

5. 38.2

98.0

6. 0.0183

7. 0.98934

8. 0.98623

9. 35.65

10. 1.33497

11. 7.77

13. 0.4769

PROBLEM SET X

1. 22.4

2. 10.764

3. -459.63

4. -40

5. 252.00

6. 62.371

PROBLEM SET XI

1. 98.22

2. 76.64

3. (a) 72.38

4. 749.35

(b) -0.08

5. (a) 0.999891

6. 26.4853

(b) 26.4853

7. 12.3865

8. 19.5370

9. 249.434

10. 0.05076

PROBLEM SET XII

1. (a) 12.0106

2. (a) 0.0084

(b) 12.0102

(b) 0.00180

(c) 12.01038

(c) 0.00235

(d) 0.0010

- | | |
|---------------|-------------|
| 3. (a) 0.0043 | 4. 12.01021 |
| (b) 0.00185 | 0.000226 |
| (c) 0.00191 | |
| Linhart's | |
| 5. 0.9665 | 6. 35.45626 |
| | 0.000613 |
| 7. 0.0000127 | 8. 12.01052 |
| | 0.00103 |

PROBLEM SET XIII

- | | |
|------------------------|------------------------|
| 2. $\frac{1}{5}_2$ | 3. $\frac{1}{1}_3$ |
| 4. $\frac{1}{2}_{13}$ | 5. (a) $\frac{1}{4}_6$ |
| | (b) $\frac{1}{5}_{16}$ |
| 6. (a) $\frac{5}{3}_6$ | 7. 84 |
| (b) $\frac{1}{1}_8$ | |
| (c) $\frac{1}{3}_6$ | |
| 8. (a) 0.205 | 9. 0.0420 |
| (b) 0.828 | |
| 10. 0.11246 | 11. (a) 0.000686 |
| | (b) 0.000163 |
| | (c) 0.000445 |
| | (d) 0.0000089 |

PROBLEM SET XIV

- | | |
|--|--|
| 1. $(6.064 \pm 0.006) \times 10^{23}$ | 2. $(1.6622 \pm 0.0017) \times 10^{-24}$ |
| 3. (a) 12.01038 ± 0.00013 | 4. 12.01024 ± 0.000043 |
| (b) 12.01021 ± 0.000058 | |
| 5. (a) 35.45626 ± 0.00016 | 6. 126.91569 ± 0.00076 |
| (b) 1.638076 ± 0.0000022 | |
| 7. Yes. | 8. 4.6176 ± 0.0020 |
| 9. $(6.5465 \pm 0.0018) \times 10^{-27}$ | 10. No |
| Ratio = 0.39 | |
| 11. (a) 7.6×10^{-4} | 12. (a) 54.94 ± 0.11 |
| (b) 5×10^{-5} | (b) 38.30 ± 0.11 |
| (c) 1.6×10^{-12} | (c) 4.894 ± 0.023 |

PROBLEM SET XV

- | | |
|--------------------------------------|-------------------------------------|
| 1. $T = 29.49M^{0.549}$ | 2. Third-degree |
| 3. $y = e^{x^2}$ | 4. $d = 0.9970 + 0.00086C$ |
| 5. $\log K = 5.55 - \frac{8,050}{T}$ | 6. (a) $y = 2.42 + 0.444x$ |
| | (b) $y = 2.52 + 0.435x$ |
| 7. $\log R = 1.700 - 0.00601C$ | 8. $\sqrt{1/\lambda} = -261 + 286N$ |

AUTHOR INDEX

A

ANDERSON, 225
ARCHIMEDES, 9, 35, 106
ARENSON, 227
AUBERT, 122
AVOGADRO, 200

B

BALL, 9, 35, 153, 227
BATES, 138
BAXTER, 138, 139, 173, 200
BECKMANN, 181, 182
BERLINER, 63
BERNOULLI, 143
BESSEL, 170, 183, 199, 222
BINGHAM, 72
BIRGE, 201, 228
BLAKE, 120
BLANCHARD, 122
BOCHER, 227
BODENSTEIN, 225
BOLTZMANN, 201
BOND, 223, 227
BOOLE, 227
BÖRNSTEIN, 115
BOWLEY, 227
BOYLE, 90
BRECKENRIDGE, 228
BRIGGS, 9
BRODHUN, 122
BURGESS, G. K., 101, 108
BURGESS, J., 232

C

CAJORI, 9, 20, 228
CAMPBELL, 53, 202, 228
CARDAN, 50
CHAN, 139, 173
CHAPPEL, 9, 10, 228

CHARLES, 90
CHAUVENET, 189, 190, 200, 201
COMTE, 30, 198
CORK, 101
COTES, 165
CRUMPLER, 226

D

DANIELS, 58, 179, 228
DARWIN, 198
DASHIELL, 120
DE FERMAT, 143
DEMING, 228
DE MORGAN, 89

E

EGAN, 75
EINSTEIN, 198
EUCLID, 9

F

FALES, 106, 228
FARADAY, 200
FINE, 46, 228
FISCHER, 229
FOWLE, 228
FRY, 228
FURMAN, 116

G

GALE, 159, 228
GAUSS, 106, 123, 165, 166
GOODWIN, 31, 211, 228
GUNTER, 21

H

HALE, 138, 139, 173, 200
HALL, 107
HALLOCK, 85, 228
HAMMOND, 36
HECHT, 122
HODGMAN, 80, 99, 228
HOFFMAN, 228

HOLMAN, 31
 HOLTZCLAW, 228
 HÖNIGSCHMID, 139, 173
 HOPKINS, 103
 HORN, 120
 HURLEY, 138, 173, 200
 HUYGENS, 143

J

JACKSON, 72
 JEFFREYS, 194, 195

K

KAYE, 228
 KELVIN, 1
 KEMP, 75
 KENNY, 106, 228
 KEYSER, 82
 KOENIG, 122
 KOLTHOFF, 106, 111, 116, 228
 KOPP, 62
 KRAUSKOPF, 225

L

LABY, 228
 LACEY, 228
 LAGRANGE, 73, 75, 76, 77, 80, 165
 LANDOLT, 115
 LANGE, 228
 LAPLACE, 143, 165, 166
 LARSON, 139
 LE CHÂTELIER, 101
 LEGENDRE, 166, 217
 LEIBNITZ, 153
 LEWIS, 64
 LINGANE, 139
 LINHART, 138
 LIVINGSTON, 228
 LOVITT, 228
 LUNDELL, 228

M

MANNHEIM, 21
 MAXWELL, 27
 MAY, 63
 MELLOR, 179, 198, 228

MENDELYEEV, 226
 MENZIES, 80, 81
 MICHELSON, 83
 MILLER, 109
 MILLIKAN, 27
 MILNE-THOMPSON, 229
 MÖLLER, 225
 MOSELEY, 226
 MURCHISON, 122

N

NAPIER, 8, 9
 NERNST, 26
 NEWTON, 58, 60, 61, 62, 63, 66, 67, 70,
 71, 73, 75, 77, 79, 153, 198
 NICOLE, 67

O

OGLESBY, 50
 OSBORNE, 229
 OUGHTRED, 20, 21

P

PALMER, 229
 PARTINGTON, 179, 229
 PASCAL, 143
 PETER, 170, 171, 199
 PIERCE, 189, 190, 201
 PIERI, 82
 PLANCK, 27, 35, 201
 PLATO, 96

R

RAE, 229
 RAMSAY, 188
 RANDALL, 64
 RAYLEIGH, 186, 188
 REED, 229
 REILLY, 229
 RICHARDS, J. W., 95
 RICHARDS, T. W., 103
 RICHTER, 35
 ROBERTS, 51
 ROBERTSON, 21
 ROBINSON, 135, 207, 222, 229

ROGERS, 103
ROGET, 21
ROSCOE, 52
ROTH, 115
ROWLAN, 101, 171
RYDBERG, 201

S

SANDELL, 106, 111, 228
SATTERLY, 229
SCHEEL, 79, 115
SCHORLEMMER, 52
SCOTT, 138, 173, 200
SHANKS, 35
SHAW, 229
SHERWOOD, 229
SIMPSON, 165
SMILES, 62
SOKOLNIKOFF, E. S., 46, 229
SOKOLNIKOFF, I. S., 46, 229
STARCK, 225
STEFAN, 201
STEFFENSEN, 229
STRIEBEL, 139, 173
SULLIVAN, 82
SWIFT, 89

T

TALMUD, 141
TAYLOR, B., 67, 168
TAYLOR, H. A., 63
TAYLOR, H. S., 63
THIESEN, 79
THOMPSON, J. E., 229
THOMPSON, S. P., 229

TREADWELL, 107
TUTTLE, 229

V

VAN CEULEN, 35
VAN DER WAALS, 56, 90
VLACQ, 9
VON VEGA, 229
VOSBURG, 138

W

WADE, 85, 228
WALKER, 224
WALTENBERG, 108
WASHBURN, 229
WATKEYS, 159, 228
WEATHERILL, 103
WEBER, 120, 121, 122
WHEELER, 229
WHITEHEAD, A. N., 125
WHITEHEAD, T. N., 229
WHITTAKER, 135, 207, 222, 229
WIEN, 201
WIRSING, 120
WOOD, 101
WOODWARD, 229

Y

YOE, 109, 120, 226
YOUNG, 61

Z

ZINN, 103

SUBJECT INDEX

A

- Accuracy, 83, 95
- Activity, 88, 89
- Addition of exponential numbers, 7
- Adsorption, surface, 110
- Analysis, gravimetric, 107
 - errors in, 107-111
- Answers, table of, 234
- Antilogs, 19
- Approximations, 198
 - by neglecting small added or subtracted quantities, 53
 - graphical, 55
 - Newton's method, 58
 - successive, 54
- Arabic mathematicians, 1
- Area, fundamental standard of, 86
- Arithmetic mean, 126
 - short-cut method for, 128
- Atomic weights, 88
 - table of, 233
- Average deviation, 131
- Averages, 93, 125
 - statistically weighted, 191

B

- Barometer, 101
- Bessel's formula, 170, 183, 199, 222
- Bibliography, 227-229
- Binomial expansion, 152
- Briggsian logs, 20

C

- Calculus of finite differences, 67
- Calibration, of thermometer, 98
 - of volumetric ware, 111
 - of weights, 103
- Calorie, 87
- Characteristic, 10

- Class interval, 136
- Coin tossing, 150
 - frequency distribution in, 154
- Combinations, 148
- Commutative law, 69
- Computation, methods, 1
 - rules for, 31
 - stoichiometrical, 36
- Congruity in measurement, 180
- Constant errors, 94
- Constant increment, 71
- Constants, evaluation of, 210
 - graphical method, 211
 - interpolation formulas for, 211
 - method of averages for, 215
 - method of least squares for, 217
 - simultaneous equations for, 214
- Coprecipitation, 110
- Corrigible errors, 96, 97, 122
- Cubic centimeter, 86
- Cubic equation, 50
- Curve fitting, 202
- Curve of error, 159
 - properties of, 162
- Curves, distribution, 142

D

- Density, 86
- Determinants, 41
 - higher order, 43
 - solving simultaneous equations by, 43, 46
- Deviation, 93, 127
 - average, 131
 - standard, 131
 - short-cut method for, 132
- Differences, 67
 - finite, 203
 - statistically reliable, 184
- Digit factor, 2

- Digits, 30
 - significant, 30
 - Dispersion, 130
 - coefficient of, 130
 - in a binomial distribution, 155
 - Distribution, 134
 - binomial, 153
 - average and dispersion in, 155
 - curves, 142
 - frequency in coin tossing, 154
 - of errors, 141
 - Distributive law, 69
 - Double salt, 111
- E
- Electrode potential, 88
 - Elimination, 39
 - general method, 40
 - Energy, 87
 - Equations, algebraic solution of, 36
 - approximate methods of solving, 53
 - choice of, 203
 - degree of, 38
 - higher degree, 50
 - linear, 36
 - simultaneous, 38
 - quadratic, 48
 - stoichiometrical, 36
 - simultaneous, 47
 - Error function, 167
 - Error integral, 167
 - table of, 232
 - Errors, 93
 - classification of, 96
 - constant, 94
 - corrigible, 96, 97
 - curve of, 159
 - distribution of, 141
 - in gravimetric analysis, 107-111
 - of estimation, 119
 - personal, 119
 - random, 96, 118, 122
 - residual, 118
 - theory of, 141
 - verification of "law" of, 171
 - Etymology, 1, 9
 - Evaluation of constants, 210
 - graphical method, 211
 - interpolation formulas for, 211
 - method of averages for, 215
 - method of least squares for, 217
 - simultaneous equations for, 214
 - Events, complex, 146
 - equally likely, 144
 - independent, 146
 - mutually exclusive, 146
 - related, 146
 - Evolution, 6
 - Expansion, binomial, 152
 - Exponent summation law, 69
 - Exponential factor, 2
 - Exponential numbers, 2
 - addition of, 7
 - division of, 4
 - multiplication of, 4
 - powers of, 6
 - roots of, 6
 - subtraction of, 7
 - Exponents, 10
 - fractional, 10
 - law of algebraic summation of, 4, 14, 16
 - Extrapolation, 77
 - graphical, 78
 - interpolation formulas used in, 78
 - other methods of, 79
- F
- Factor, digit, 2
 - exponential, 2
 - Factorial n , 148
 - Finite differences, 203
 - Fractional exponents, 10
 - Fugacity, 89
- G
- Gas law, 90
 - Gaussian theory of errors, 165, 166, 174, 184, 194, 217
 - Glass fatigue, 100
 - Graphical method, 205, 211

Gravimetric analysis, 107
 errors in, 107-111
 Gunter's rule, 21

H

Heat quantity, 87

I

Increment, constant, 71
 variable, 73
 Interpolation, 13
 formulas, 211
 inverse non-linear, 77
 linear, 64
 non-linear, 66
 Involution, 6

K

Kilogram, 84

L

Lagrange's interpolation formula, 75
 "Law" of errors, verification of, 171
 Least squares, principle of, 166
 Length, fundamental standard of, 83
 Linear equations, 36
 Liter, 86
 Logs (Logarithms), 8
 common (Briggsian), 20
 division by use of, 14
 extracting the root of a number by
 use of, 16
 invention of, 8
 multiplication and division by use
 of, 15
 multiplication by use of, 12
 natural (Napierian), 19
 negative, 16
 raising a number to a power by use
 of, 15
 slide rule, 26
 table of, 230-231

M

Mantissa, 10
 Mass, fundamental standard of, 84
 Mean, arithmetic, 126
 weighted, 127
 Measurements, congruity in, 116
 general discussion of, 89
 obtaining congruity in, 180
 statistical interpretation of, 174
 theory of, 82
 Median, 126
 Meter, 83
 Method, of averages, 215
 of least squares, 217
 Milliliter, 86
 Mixed crystals, 110
 Mode, 125

N

Napierian logs, 19
 Napier's "radix", 9
 Negative logs, 16
 Newton's approximation method, 58
 Newton's difference formula, 70
 Newton's divided difference formula,
 73

O

Observations, 91
 limited number of, 194
 rejection of suspected, 188
 Occlusion, 110
 Operators, 69
 Oughtred's slide rule, 20

P

Parallax, 114
 Permutations, 147
 Personal errors, 119, 122
 Peter's formula, 170, 171, 199
 pH, 18
 calculation, 18
 problems, 20, 28, 52

- Polynomials, 66
 Post-precipitation, 111
 Potential, electrode, 88
 Powers, by use of logs, 15
 of exponential numbers, 6
 on slide rule, 25
 Precipitate, colloidal dispersion of, 110
 contamination of, 110
 Precipitation, post-, 111
 simultaneous, 110
 Precision, 94
 measure of, 169
 numerical statement of, 195
 Pressure, corrections in measurements of, 101
 unit of, 87
 Probability, mathematical, 143
 empirical, 150
 historical statement of, 143
 theorems, 143
 limitations of, 150
 Probability integral, 167
 Probability theorems, 143
 limitations of, 150
 Probable error, 169
 in complex functions, 178
 of arithmetic mean, 183
 of a computed result, 175
 of a product, 178
 of a quotient, 178
 of a sum or difference, 177
 of exponential functions, 178
 of logs, 179
 of powers and roots, 178
 Problems, Set I (exponential numbers), 3
 Set II (exponential numbers), 5
 Set III (exponential numbers), 8
 Set IV (logs), 20
 Set V (slide rule), 26
 Set VI (significant figures), 34
 Set VII (algebraic), 51
 Set VIII (approximations), 63
 Set IX (interpolation and extrapolation), 79
 Set X (units), 95
 Problems, Set XI (calibration, standardization, etc.), 123
 Set XII (statistical), 138
 Set XIII (errors, probability, etc.), 173
 Set XIV (statistical interpretation), 200
 Set XV (curve fitting, evaluation of constants, etc.), 224

 Q
 Quadratic equations, 48

 R
 Random errors, 96, 118, 122
 Range, 130
 Ratio, 1, 9
 Rejection of suspected observations, 188
 Residual errors, 118
 Roman numerals, 1
 Root mean square, 126
 Roots, by use of logs, 16
 of exponential numbers, 6
 on slide rule, 25
 Rules for computation, 31

 S
 Significant figures, 30, 197
 Simultaneous equations, 214
 solving by determinants, 43, 46
 Simultaneous precipitation, 110
 Slide rule, 20
 circular, 21
 division by, 24
 history of, 20
 K and E (fig.), 23
 log-log, 21
 logs by, 26
 multiplication by, 24
 Oughtred's, 20
 powers by, 25
 problems, 26, 29
 roots by, 25
 Smoothing of data, 207

Specific gravity, 87
 Standard deviation, 131
 short-cut method for, 132
 Standards, 82
 derived, 86
 density, 86
 energy, 87
 heat quantity, 87
 pressure, 87
 specific gravity, 87
 fundamental, 83
 area, 86
 length, 83
 mass, 84
 primary, 83
 secondary, 85
 temperature, 84
 time, 84
 volume, 86
 relative, 87
 Statistical interpretation of measurements, 174
 Statistical methods, 125
 Statistically reliable difference, 184
 interpretation of, 186
 Statistically weighted averages, 191
 Stoichiometrical computations, 36
 Substituted variables, 78, 209
 Subtraction of exponential numbers, 7
 Surface adsorption, 110

T

Table, of answers, 234
 of atomic weights, 233
 of error integral, 232
 of logs, 230-231

Taylor's theorem, 168
 Temperature, fundamental standard of, 84
 measurement of, 97
 electrical, 100
 Thermometer, Beckmann, 101, 182
 calibration of, 98
 emergent stem correction, 99
 gas, 97
 glass fatigue of, 100
 mercury-in-glass, 97
 Time, fundamental standard of, 84
 Tossing, coin, 150
 frequency distribution in, 154

V

van der Waals' equation, 56, 91
 Variable increment, 73
 Variables, substituted, 78, 209
 Volume, fundamental standard of, 86
 measurement of, 111
 temperature corrections of, 115

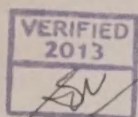
W

Weber's law, 120, 122
 Weighing, 103
 corrections in, 105, 106
 Weighted average, 191
 Weights, calibration of, 103

Y

Young's equation, 61

~~NL~~ 318789 ~~NL~~ 16-4-97



CFTRI-MYSORE



229

Chemical computa.

Acc. No. 229

2 N47

PLER (TB).

Compu-

s.

INTERNATIONAL ATOMIC WEIGHTS

1943

Reprinted from the *Journal of the American Chemical Society*

	Sym- bol	Atomic Number	Atomic Weight		Sym- bol	Atomic Number	Atomic Weight
Aluminum...	Al	13	26.97	Molybdenum..	Mo	42	95.95
Antimony....	Sb	51	121.76	Neodymium..	Nd	60	144.27
Argon.....	A	18	39.944	Neon.....	Ne	10	20.183
Arsenic.....	As	33	74.91	Nickel.....	Ni	28	58.69
Barium.....	Ba	56	137.36	Nitrogen.....	N	7	14.008
Beryllium....	Be	4	9.02	Osmium.....	Os	76	190.2
Bismuth.....	Bi	83	209.00	Oxy.....	O	8	16.0000
Boron.....	B	5	10.82	Palladium	Pd	46	106.7
Bromine.....	Br	35	79.916	Phosphorus	P	15	30.98
Cadmium....	Cd	48	112.41	Platinum	Pt	78	195.23
Calcium.....	Ca	20	40.08	Po.....	K	19	39.096
Carbon.....	C	6	12.010	Pr.....	Pr	59	140.92
Cerium.....	Ce	58	140.12	P.....	Pa	91	231
Cesium.....	Cs	55	132.91	I.....	Ra	88	226.05
Chlorine.....	Cl	17	35.453		Rn	86	222
Chromium... Cr	24				Re	75	186.31
Cobalt.....	Co	27	58.933		Rh	45	102.91
Columbium..	Cb	41	92.906		Rb	37	85.48
Copper.....	Cu	29	63.546		Ru	44	101.7
Dysprosium..	Dy	64	162.50		Sm	62	150.43
Erbium.....	Er	68	167.26		Sc	21	45.10
Europium....	Eu	63	151.96		Se	34	78.96
Fluorine.....	F	9	18.998		Si	14	28.06
Gadolinium..	Gd	64	157.25	Silicon.....	Si	14	28.06
Gallium.....	Ga	31	69.723	Silver.....	Ag	47	107.880
Germanium..	Ge	32	72.64	Sodium.....	Na	11	22.997
Gold.....	Au	79	196.967	Strontium...	Sr	38	87.63
Hafnium.....	Hf	72	178.49	Sulfur.....	S	16	32.06
Helium.....	He	2	4.003	Tantalum....	Ta	73	180.88
Holmium....	Ho	67	164.93	Tellurium....	Te	52	127.61
Hydrogen....	H	1	1.0080	Terbium.....	Tb	65	159.2
Indium.....	In	49	114.71	Thallium....	Tl	81	204.39
Iodine.....	I	53	126.905	Thorium.....	Th	90	232.12
Iridium.....	Ir	77	192.22	Thulium....	Tm	69	169.4
Iron.....	Fe	26	55.847	Tin.....	Sn	50	118.70
Krypton.....	Kr	36	83.74	Titanium....	Ti	22	47.90
Lanthanum..	La	57	138.905	Tungsten....	W	74	183.92
Lead.....	Pb	82	207.19	Uranium....	U	92	238.07
Lithium.....	Li	3	6.941	Vanadium....	V	23	50.95
Lutecium....	Lu	71	174.967	Xenon.....	Xe	54	131.3
Magnesium..	Mg	12	24.305	Ytterbium...	Yb	70	173.04
Manganese..	Mn	25	54.938	Yttrium.....	Y	39	88.92
Mercury.....	Hg	80	200.59	Zinc.....	Zn	30	65.38
				Zirconium....	Zr	40	91.22

